

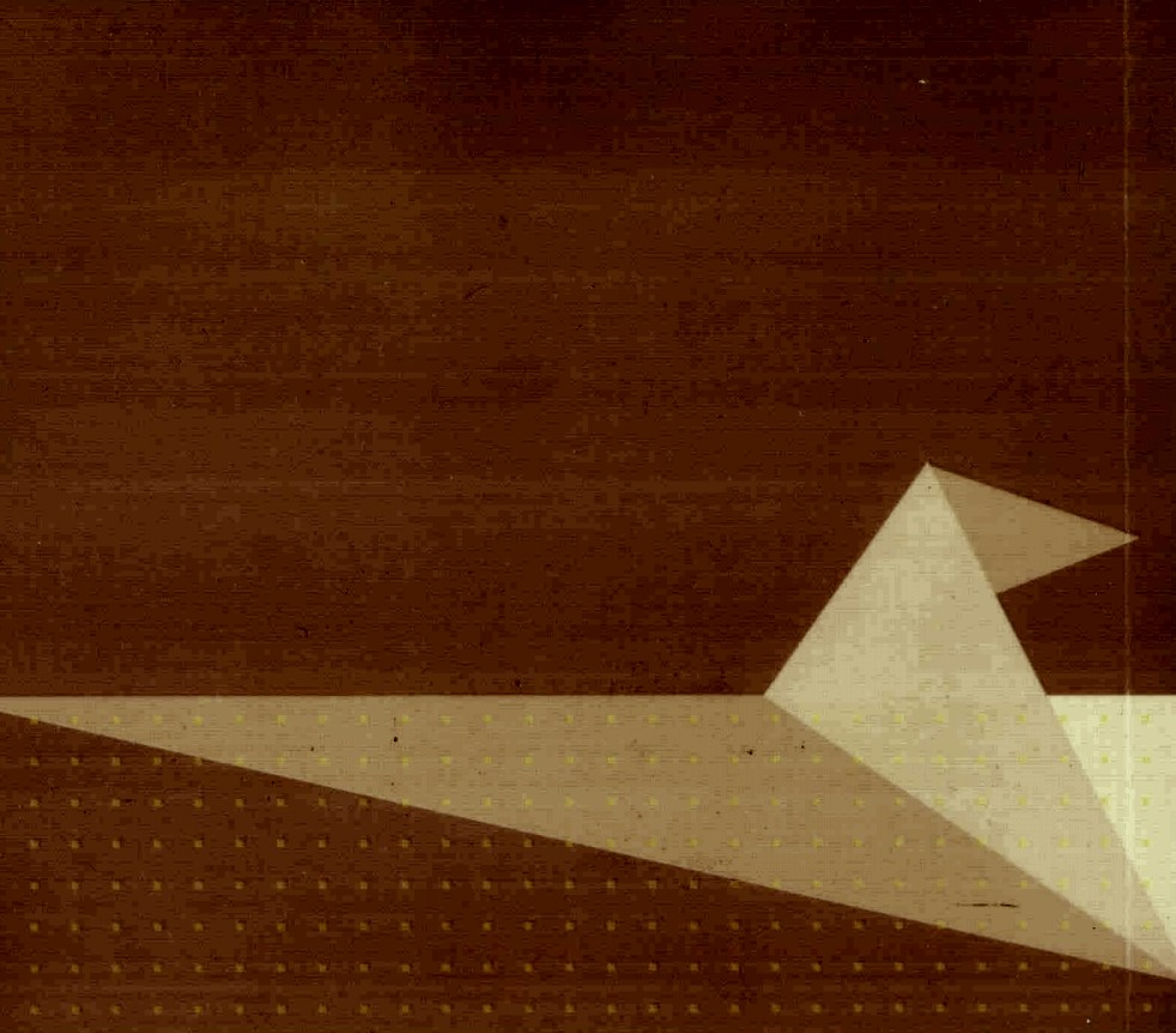
VISIT...

LANZAROTE
Caliente.COM

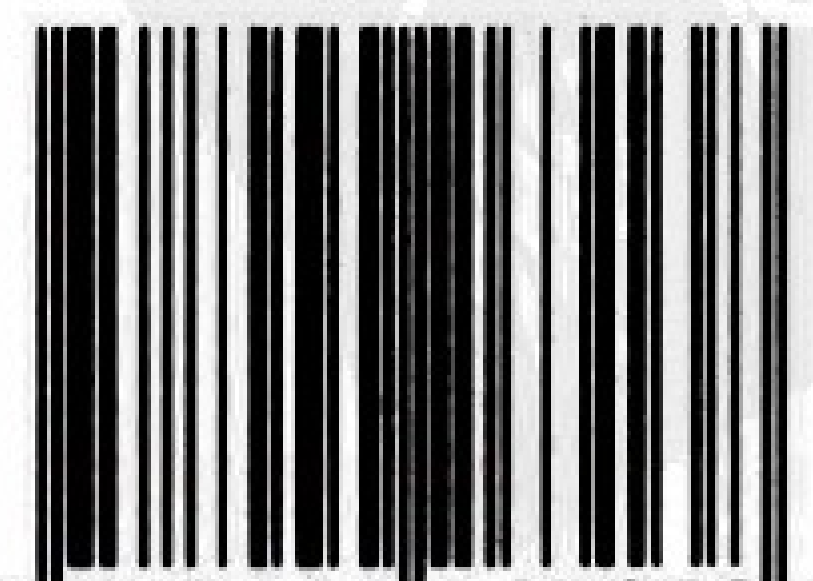
当代语言学

□ 陈保亚 著

高等教育出版社



ISBN 978-7-04-026313-8



9 787040 263138 >

定价 30.10 元

当代语言学

陈保亚 著

高等教育出版社

图书在版编目 (CIP) 数据

当代语言学 / 陈保亚著. —北京: 高等教育出版社,
2009. 7

ISBN 978 - 7 - 04 - 026313 - 8

I. 当… II. 陈… III. 语言学 - 高等学校 - 教材
IV. H0

中国版本图书馆 CIP 数据核字 (2009) 第 095409 号

首席策划 袁晓波 策划编辑 于晓宁 责任编辑 吴 军
封面设计 王凌波 版式设计 范晓红 责任校对 王效珍
责任印制 陈伟光

出版发行	高等教育出版社	购书热线	010 - 58581118
社 址	北京市西城区德外大街 4 号	咨询电话	400 - 810 - 0598
邮政编码	100120	网 址	http://www.hep.edu.cn
总 机	010 - 58581000		http://www.hep.com.cn
		网上订购	http://www.landracom.com
经 销	蓝色畅想图书发行有限公司		http://www.landracom.cn
印 刷	北京奥鑫印刷厂	畅想教育	http://www.widedu.com
开 本	850×1168 1/32	版 次	2009 年 7 月第 1 版
印 张	19.875	印 次	2009 年 7 月第 1 次印刷
字 数	490 000	定 价	30.10 元

本书如有缺页、倒页、脱页等质量问题, 请到所购图书销售部门联系调换
版权所有 侵权必究

物料号 26313 - 00

内容简介

本书分析了当代语言学的主要共时理论模型、观点和操作方法,对一些主要概念作了清理,对存在的某些问题提出了自己的初步解决方案。全书重点围绕从结构语言学到转换生成语法的转型这条主线展开。作者重新讨论了结构语言学分析手续的合理部分,对结构语言学发现程序在田野调查工作中的重要价值作了进一步论证并讨论了结构语言学的线性原则在语言生成过程中的简单性、必要性和存在的问题,分析了转换生成语法及其相关理论在解决这些问题上的具体方案和重要价值,进一步论证了语法功能、语义功能在分析中的必要性和重要性。作者把句子的生成过程分成生成程序和生成条件,指出转换生成语法在生成程序上做了很多重要工作,取得了重大进展,但在生成条件方面做得相当不够,所以生成合法句子的理想远远没有达到。作者最后认为今后的工作应该更加注重生成条件的研究,即线性组合规则和非线性转换规则中平行周遍条件的研究。

作者十多年来一直在北京大学中文系给研究生讲授当代语言学。本书是对讲稿主要内容的整理,是作者多年教学、研究和思考的一个总结,可作为大学语言类专业的选修课教材或参考书。

郑重声明

高等教育出版社依法对本书享有专有出版权。任何未经许可的复制、销售行为均违反《中华人民共和国著作权法》，其行为人将承担相应的民事责任和行政责任，构成犯罪的，将被依法追究刑事责任。为了维护市场秩序，保护读者的合法权益，避免读者误用盗版书造成不良后果，我社将配合行政执法部门和司法机关对违法犯罪的单位和个人给予严厉打击。社会各界人士如发现上述侵权行为，希望及时举报，本社将奖励举报有功人员。

反盗版举报电话：(010) 58581897/58581896/58581879

反盗版举报传真：(010) 82086060

E - mail：dd@hep.com.cn

通信地址：北京市西城区德外大街4号

高等教育出版社打击盗版办公室

邮 编：100120

购书请拨打电话：(010)58581118

目 录

引言	1
1. 言语链中语素的切分	5
1.1 Saussure 线条性的实质	5
1.2 双项对比	6
1.3 对双项对比原则的进一步限制	8
1.4 对比程序和可对比程度	11
思考问题	13
2. 线性序列的生成程序	14
2.1 替换与连接	14
2.2 扩充与扩展	15
2.3 叠加	21
思考问题	31
3. 单位和规则	32
3.1 词作为句法单位遇到的问题	32
3.2 单位及组合的可推导性	34
3.3 平行周遍原则与生成规则	51
3.4 平行的性质	54
3.5 周遍的性质	58
3.6 多个语素序列的判定	62
3.7 语素库、语符库和词库	65

3.8	派生、转换和语符库的关系	68
3.9	短语的语类推导	73
	思考问题	79
4.	替换、叠加和成分线性特征	80
4.1	自动机理论与自然语言的成分线性特征	80
4.2	替换的有向性与有限状态语法的左线性特征	92
4.3	有限状态语法的根本属性	96
4.4	替换、叠加与短语结构语法	100
4.4.1	短语结构语法($[\Sigma, F]$)的基本结构	101
4.4.2	短语结构语法的本质	104
4.5	有限状态语法和语类规约模型	110
	思考问题	120
5.	线性原则的不充分性和转换的必要性	121
5.1	Harris 的线性化工作	121
5.2	转换问题的提出	122
5.3	Chomsky 关于转换存在的理由	126
5.4	转换规则的必要性	131
5.4.1	英语助动词和转换	133
5.4.2	主动和被动的关系	134
5.4.3	线性组合程序和不连续直接成分	137
5.4.4	并合与转换	140
5.4.5	非线性组合	141
	思考问题	144
6.	音系分析程序的非线性化	145
6.1	自主音系学的不充分性	151
6.2	语素基本音形与变音	153

6.3	音系分析条件:语素基本音形	155
6.4	从语素音位到语音规则和语素基本音形	162
6.5	核心语音单位	172
6.6	提取核心语音单位以语素基本音形的长度为必要条件 ...	180
6.7	提取核心语音单位的语境不能短于语素音形	191
6.8	语素音形的一致性和独立语素音形的确定	196
6.9	平行周遍条件与语音规则	207
6.10	可延展性与音节的判定	214
6.11	两类语音组合规则	237
6.11.1	音段组合成音节	237
6.11.2	音节和音节的组合	239
6.12	对立矩阵与音位的提取程序	244
6.12.1	对立研究的必要性	246
6.12.2	元音音位的提取	253
6.12.3	辅音音位的提取	267
6.12.4	音子组合数与对立环境	271
	思考问题	280
7.	语义组合关系的线性分析	281
7.1	组合规则的语义解释	281
7.2	次范畴化、词库和平行周遍条件	284
7.2.1	名词的次范畴化和动词的选择规则	285
7.2.2	动词次范畴化	289
7.2.3	次范畴化和平行周遍条件	292
7.3	深层结构及其线性特征	299
7.3.1	深层结构的存在依据	303
7.3.2	深层结构的确定	312
7.3.3	深层结构的线性顺序	314
	思考问题	317

8. 语义组合关系的非线性化及其功能分析	318
8.1 格语法的生成机制	319
8.2 语义格和线性深层结构的关系	329
8.3 语义格的初始性	334
8.4 词库中动词语义格的确定和直接关联原则	344
8.5 动词的格框架和标注	355
8.6 语义结构变体标注	366
8.7 论元结构和次范畴结构的相互独立性	376
8.8 文本中语义格的识别	385
8.9 多项式的判定问题	389
思考问题	396
9. 语义结构的再次非线性化:语义表达式	397
9.1 深层结构在语义解释上遇到的问题	397
9.2 语义表达式及其深度	399
9.2.1 表达深度 1:原子命题的组合差别	401
9.2.2 表达深度 2:词汇分解	406
9.3 深层结构中的词汇插入问题	408
9.4 元语言问题	411
9.4.1 自然语言是最初始的元语言	412
9.4.2 自然语言自身解释的内部循环	414
9.4.3 获得元语言的可能性:核心语符和核心规则	418
9.4.4 核心规则的建构分析	421
9.5 语义表达式的初始性	429
思考问题	433
10. 语义解释层面	434
10.1 语义格框架和投射原则	434

10.2 语义结构的语义和转换的语义	439
10.3 平行转换中的语义差别	449
思考问题	455
11. 结构的阶与递归性	456
11.1 X 杠理论	457
11.2 限定语和补足语层次的经验依据	469
11.3 附加成分的再分类	474
11.4 两种递归方式	480
11.5 X 杠的递归性	482
11.6 短语结构的阶和语杠理论	489
11.7 各种语类的阶和杠	508
11.8 X 杠理论的核心问题	511
思考问题	515
12. 句法结构关系的初始性	516
12.1 支配和约束	517
12.2 支配概念的扩展和长距离约束的传递条件	521
12.3 支配与结构关系	529
12.3.1 c-command	529
12.3.2 m-command	536
12.3.3 支配的本质和作用	540
12.4 句法结构关系的初始性	550
思考问题	557
结语:从程序到条件	558
符号解释	567
人名译名对照外文排序	569
人名译名对照汉语排序	571

主要概念译名对照外文排序	573
主要概念译名对照汉文排序	582
索引	591
参考书目	599
后记	624

引言

话语的理解和生成需要哪些初始概念和程序？这是本书要集中讨论的第一个基本问题。

结构语言学以前的传统语法比较重视单位的功能，那时有两个层面的功能概念是相当重要的，一个是主语、宾语等句子成分的概念，另一个是施事、受事或逻辑主语、逻辑宾语的概念，相当于后来的论元角色。20 世纪初，从 Saussure 开始，不同学派对待功能的态度有了分歧。Saussure 从线性原则出发，强调聚合关系和组合关系（句段关系）。他所说的这两种关系实际上是互为因果的，即组合关系从根本上说就是话语线条中可以相互替换的聚合类所形成的组合关系。Saussure 没有讨论句子成分和逻辑成分这两种功能概念。布拉格学派很重视功能的研究，尤其是对话题述题有很多讨论。美国结构语言学则更多地采取了一种经验观察的态度。Bloomfield (1926) 只是在词的分布问题上谈到功能，他所说的功能常常是指语类，传统语法中的句子成分、语义角色、话题等都不太考虑。从这种意义上说，美国结构语言学和 Saussure 的研究思路更为一致，和布拉格学派有所不同。Bloomfield (1933) 只用了向心结构和离心结构的概念来讨论单位之间的结构关系，其他功能概念都不涉及。1946 年，Harris 发表了 *From morpheme to utterance*，认为从语素的分布开始，可以生成全部的话语，所有的功能关系，包括语法结构关系、语义结构关系、表达关系，都不再考虑，这时我们说 Harris 已经完全和功能主义分道扬镳，走向形式化的道路。Harris 的研究和 Bloch 的音系理论有共同的地方，就是希望找出一种发现程序，即通过最少的概念和原则，提取单位，给单位

分类,并通过单位的分布类说明组合关系。我们可以把这种以发现程序为研究目标的研究取向称为后结构语言学。从 1957 年开始,Chomsky 又采取了一种不同于后结构语言学的研究态度,反对发现程序,提出了评价程序的思想,开始利用假说模型来完成句子的生成。至于功能,直到 20 世纪 60 年代 Chomsky (1965)的标准理论都不加考虑。那时 Chomsky 认为功能是可以通融的。

以上转型过程显示了三种不同的研究态度,各自主要体现在功能语言学、结构语言学和转换生成语法中,由此形成了三种主要的研究范式。早期的结构语言学和功能语言学在很多方面都是相容的,只是到了后结构语言学的发现程序才开始和功能语言学分道扬镳,但在很多语言观念上基本上还是可以通融的,尤其在经验主义的态度上是一致的。从结构语言学向转换生成语法的转型是 20 世纪最为令人瞩目的转型,因为在很多语言观念上都有了根本的改变,尤其是从经验主义向唯理主义(rationalism)的转变。Chomsky(1957)的 Syntactic Structure 是根本的转折点。但是,截至目前,转换生成语法并没有完全取代结构语言学和功能语言学。在田野调查中,人们更多地用结构语言学的方法,比如单位的提取首先要考虑替换和分布等因素。在语言学习和语言交流的研究中,很多人在用功能语言学的方法。一个外国人的汉语作文可能每个句子都正确,但仍然会前言不搭后语,因为他并没有把握住语篇布局和句子的用法。在语言机制和语言认识论研究方面,更多的人偏向转换生成语法,也有吸收各方面成果的研究。但是,这三大范式的对立点在哪里,有没有内在关系,如果有,这种关系是什么?各个范式的优劣在哪里?这些方面一直没有得到很好的研究。

当后结构语言学抛弃功能概念的时候,并没有充分论证抛弃的理由。当转换生成语法反对发现程序的时候,只是说发现程序目标太高,并没有证明是否可以绕开发现程序完成对语言

的认识。如果我们要调查一个陌生语言,首先要有办法获得这个语言的单位,这就是发现程序的工作,转换生成语法绕开了这个根本问题。不过功能语言学和结构语言学本身也没有证明功能和发现程序的必要性。结构语言学尽管提出了提取单位的发现程序,但发现程序绕开规则找单位,也遇到了严重的问题。转换生成语法的很多概念和原则,还是很有价值的,但很多也都没有得到证明,这是转换生成语法遭到反对的重要原因。这一系列问题成为本书要重点讨论的问题,概括起来就是引言一开始提出的基本问题:话语的理解和生成需要哪些初始概念和程序?

有了一套初始概念和程序是否就一定能够理解和生成合法的句子?比如,如果我们有了转换规则,就可以说复合动趋式中宾语中置、前置是通过转换获得的:

- | | |
|-----|--|
| 后置: | 我从冰箱里拿出来[一瓶酒] |
| | 我从柜子里拿出来[一件衣服] |
| 中置: | 我从冰箱里拿出[一瓶酒] _i 来[t _i] |
| | 我从柜子里拿出[一件衣服] _i 来[t _i] |
| 前置: | 我从冰箱里拿[一瓶酒] _i 出来[t _i '] |
| | 我从柜子里拿[一件衣服] _i 出来[t _i '] |

构式理论可以不承认转换层面,说这里是三种不同的构式。结构语言学也可以说中置、前置句式中存在不连续直接成分。我们将会讨论这些理论模型的关系。这里的关键问题是,无论采取哪种理论模型,一旦我们加上一个“刚才”,合法性问题就开始出现:

- | | |
|-----|----------------|
| 后置: | 我刚才从冰箱里拿出来一瓶酒 |
| | 我刚才从柜子里拿出来一件衣服 |
| 中置: | 我刚才从冰箱里拿出一瓶酒来 |

我刚才从柜子里拿出一件衣服来
前置： ？我刚才从冰箱里拿一瓶酒出来
 ？我刚才从柜子里拿一件衣服出来

这是各种模型都需要解释的问题。类似的问题是什么引起的，目前缺乏统一的解释。总体而言，当前对发现程序和生成程序的研究工作已经做了很多，但对还原单位、规则以及生成合法句子的条件依据什么原则却研究得很少。这是本书要追问的另一个基本问题。

由于结构语言学和功能语法在经验论这个层面可以大致统一起来，我们的分析重点围绕从结构语言学到转换生成语法的转型这个线路展开，并在其中讨论跟转型密切相关的语法内部的功能问题。语篇功能问题本书不专门讨论。具体思路是先从线性程序开始分析，然后依次讨论句法层面、音系层面和语义层面的非线性问题以及转换的必要性，再分析语法结构功能（句法结构功能）的必要性。最后讨论单位和规则还原的必要性以及从组合程序研究向组合条件研究转向的迫切性。

1. 言语链中语素的切分

1.1 Saussure 线条性的实质

20 世纪语言研究中的一个显著趋势是从结构语言学到转换生成语法的转向,这主要是从对线性原则的认识开始的,这就把我们引向了对 Saussure 线性观念的认识。Saussure 说言语遵循线条原则。实际上不同的人对这一原则可以作出很多不同的理解,在后面的章节中我们会逐渐展开这个问题。有一点大家可以有共同的理解:单位在时间中前后一个个相继出现。不过单位是有层面的,可以是区别层面的单位,比如音位,也可以是表达层面的单位,比如语素或词。从 Saussure 关于线条性的定义我们可以归纳出两种含义,一种可以称为能指线条性。Saussure(1916,p106)指出:

能指属听觉性质,只在时间上展开,而且具有借自时间的特征:(a)它体现一个长度,(b)这长度只能在一个向度上测定:它是一条线。

Saussure 还认为:

这个原则(指线性原则)是显而易见的,但似乎为常人所忽略,无疑是因为大家觉得太简单了。然而这是一个基本原则,它的后果是数之不尽的;它的重要性与第一条规律(即任意性原则)不相上下。语言的整个机构都取决于它。

不过在讨论句段关系(组合关系)的时候, Saussure 把线条性扩展到了词或单位, 我们可以称为符号线条性(Saussure, 1916, p170—171):

一方面, 在话语中, 各个词, 由于它们是连接在一起的, 彼此结成了以语言的线条特性为基础的关系, 排除了同时发出两个要素的可能性, 这些要素一个挨一个排列在言语的链条上面。这些以长度为支柱的结合可以称为句段。所以句段总是由两个或几个连续的单位组成的……一个要素在句段中只是由于它跟前一个或后一个, 或者前后两个要素相对立才取得它的价值。

Saussure 认为确定线条中的单位是相当困难的, 他给出了一种提取单位的基本方法: 替换。沿着 Saussure 的主要思路发展出来的结构语言学, 一项重要的工作就是在线性序列中通过替换提取单位, 通过分布归纳单位的类。

1.2 双项对比

Saussure 给出了单位在线条性中的两种关系, 即聚合关系和组合关系。考虑 Saussure 的一个法语实例 défaire(解除), Saussure 认为在人们的下意识中存在着一个或几个联想系列, 其中一部分相同(Saussure, 1916, p179):

défaire	
décoller	faire
déplacer	refaire
découdre	contrefaire
etc.	etc.

Saussure 根据这种关系断定 *défaire* 是句段(组合)。如果包含 *dé-*和 *faire* 的其他形式从语言中消失,*défaire* 就不是句段而是单位了。Saussure 已经系统陈述了通过对比区分单位和组合的基本思想。

不过在不同的层面,对比和替换具有很多不同的性质,Saussure 当时并没有充分考虑。我们将逐渐展开分析。首先,线性组合中语素是一种重要的单位。仔细分析结构语言学以后的生成语法、功能语法以及其他各种语法理论,都不能绕开语素的提取问题,所以语素的提取问题有必要先加以考虑。

从更严格的意义上说,提取语素应该分成两个步骤,首先从言语链条中切分语子,然后把语子归并成语素,但这需要更复杂的表述,这里不展开。下面只讨论语素的切分,实际上是语子的切分。

Bloomfield(1933,10.1)把对比的方法更进一步具体化,最早提出了通过对比(音义同一性对比)切分语素。问题就在于怎样切分剩余语素,即 Bloomfield 所说的独一无二的成分(unique constituent)。在大多数情况下,这种方法是有效的。不过,有的言语片断只有一部分能够作音义同一性对比:

苹果	苹果
?	芒果
?	水果
.....	

“苹果”的“果”能够找到进行音义同一性对比的材料,“苹”找不到音义同一性对比的材料。根据对比原则可以把“果”看成语素,把“苹”看成剩余语素。这种情况属于我们曾经提到的单项对比(陈保亚,1999,6.1)。单项对比方法贯彻到底会出现下面的问题:

mei ⁵¹ 妹	mei ⁵¹ 妹
mu ²¹⁴ 姆	?
ma ⁵⁵ 妈	?

似乎也可以说“妹、姆、妈”在语音上有相同的部分 m-, 其意义表示“女性”, 满足音义同一性对比。我们是否需要承认“mei⁵¹ 妹”中的 m- 是表示“女性”的语素, 而把 -ei⁵¹ 看成是表示“晚辈或年少”的剩余语素? 为了解决这里的困难, 我们曾经从操作上给出了剩余语素切分的双项对比原则: 如果 A 是一个剩余语素, 与 A 组合的另一个片断必须有资格出现在其他可以进行双向对比的言语片断中。

根据这种限制, “果”可以出现在“水果”这样的片断中, 并且“水果”可以进行双向对比:

水果	水果
水车	鲜果
水壶	酸果
.....	

所以“苹”是剩余语素。“mei⁵¹ 妹”中的 m- 不能出现在可以进行双向对比的实例中, 所以不是表示“女性”的语素, -ei⁵¹ 也不是剩余语素。

1.3 对双向对比原则的进一步限制

不过双向对比原则会遇到问题。英语 when 中的 wh- 是可对比的:

when	when
where	?
which	?

why	?
what	?
.....	

而且包含 wh-的 where 可以作双项对比：

where	where
what	here
why	there
.....	

where 的两个部分 wh-和-ere 都可以进行对比,wh-似乎有表示疑问的意思,-ere 似乎有表示处所的意思,按照我们早期给出的双项对比原则,是否需要承认 when 中的 wh-是语素,-en 是剩余语素? 更一般地说,如果仅仅根据双项对比切分语素,我们就得承认存在下面的语素:

wh- (<u>wh</u> y)	表示疑问
-y (wh <u>y</u>)	表示原因
-at (wh <u>at</u>)	表示事物
-en (wh <u>en</u>)	表示时间
-ere (wh <u>ere</u>)	表示处所
th- (<u>th</u> ere)	表示远指
h- (<u>h</u> ere)	表示近指

where 的对比说明,双项对比原则可以把“妹”中的 m-排除在语素之外,但还不能把 wh-排除在语素之外。汉语中“骰子”的对比性质表面上看起来和 where 类似。

下面考虑怎样进一步限制双项对比原则。我们再来深入分析一下对比的性质,前面提到的“果”,不仅能够出现在双项对比

的言语片断中,而且“果”无论处在左项(前项)或右项(后项),包含“果”的言语片断有的也可以进行双项对比:

双项对比实例(“果”为右项)		双项对比实例(“果”为左项)	
鲜果	鲜果	果树	果树
鲜鱼	水果	果核	枣树
鲜菜	酸果	果仁	桃树
.....			

我们说“果”能出现在不定位双项对比的言语片断中。

“第一”中的“第”,对比性质和“果”不完全相同,“第”所出现的言语片断尽管能进行双项对比,但“第”却是定位的:

双项对比实例(“第”为右项)		双项对比实例(“第”为左项)	
		第一	第一
		第三	初一
		第四	十一
.....			

我们说“第”只能出现在定位双项对比言语片断中。不过,与“第”组合的言语片断“一”可以出现在不定位双项对比的言语片断中:

双项对比实例(“一”为左项)		双项对比实例(“一”为右项)	
一个	一个	第一	第一
一顿	两个	第三	初一
一件	三个	第四	十一
.....			

这和前面谈到的 wh-的情况有所不同。如果假定 wh-有表示疑问的语义,-ere 有表示处所的语义,则 wh-只能出现在言语片断的左边,-ere 只能出现在言语片断的右边,两者都只能进行定位同一性对比。也就是说,不仅 wh-是定位的,与 wh-组合的另一

个片断-ere 也是定位的。这说明英语中的 wh-和汉语的“果”、“第-”都不同。如果我们要求与剩余语素组合的另一个非剩余语素必须可以出现在不定位双项对比的言语片断中,我们就可以把“苹果”中的“苹”分析成剩余语素。“where”中的 wh-由于不能出现在不定位双项对比的片断中,不是语素,与 wh-组合的-ere 也不会被当成剩余语素。“妹”中的 m-由于不能出现在不定位双项对比的片断中,也不会被当成语素,与 m-组合的-ei⁵¹ 也不会被当成剩余语素。

但是这样的限制似乎太强大了,需要略作调整。因为按照这种限制,“骰子”的“-子”由于不能出现在不定位双项对比的片断中,“骰-”就不会被当做剩余语素。不过,“-子”尽管不能出现在不定位双项对比的片断中,但“-子”可以和“桌”等很多言语片断组合,不仅“桌子”是可以进行双项对比的,而且“桌”可以出现在不定位双项对比的片断中,比如“方桌”、“桌布”等,所以“-子”仍然可以通过间接的不定位双项对比被断定成一个语素,于是“骰”也是剩余语素。

根据以上分析,我们可以给剩余语素对比法一个操作方法上的不定位双项对比限制(严式限制):

与剩余语素 X 组合的另一个成分 Y 不仅是可对比的,并且至少应该具有以下两种性质之一:

1. Y 能出现在不定位双项对比的言语片断中。
2. 与 Y 组合的 Z 能出现在不定位双项对比的言语片断中。

1.4 对比程序和可对比程度

根据以上分析,我们给出一个提取语素的一般程序:

1. 如果 A 能出现在不定位双项对比的言语片断中,A 就是

不定位言语片断。

2. 如果 B 只能出现在定位对比的言语片断中,但同时能和
不定位言语片断 A 组合,B 就是定位言语片断。

3. 如果 C 能和 A 或 B 组合,C 也是言语片断。

4. 如果通过以上方式得到的一个言语片断 D,继续分成两部分以后,没有任何一部分具有以上 A 或 B 或 C 的对比性质,
D 就是语素。

5. 剩余语素是只能和 A 或 B 组合但自身不能进行音义同一性对比的语素。

以上对比程序也应该适合不连续直接成分的对比,但在表述上需要进行调整。

其实以上程序的制定意味着我们在对比的不同性质面前采取了一种特定的态度。对于一个言语片断来说,从可对比程度的强弱看,可以有下面几种情况:

1. X 可以出现在不定位双项对比的言语片断中,比如“水果、果树”中的“果”。

2. X 只能进行定位对比,但 X 可以和不定位置言语片断组合,比如“桌子”中的“-子”。

3. X 只能进行定位对比,并且和 X 组合的片断也只能进行定位对比。比如英语 where 中的 wh-和-ere。

4. X 只能进行定位对比,和 X 组合的片断不能进行对比,如汉语“妹”中的 m-。

如果我们只承认 1、2 两种对比中的言语片断是语素,与此相应的剩余语素也只有两种,一种剩余语素是和不定位置语素组合,如“苹果”中的“苹”,一种剩余语素是和定位语素组合,如“骰子”中的“骰”。

如果我们把可对比的标准放宽到第 3 个条件,那么英语中的 wh-也是语素,与此相应的剩余语素又多出一个类型,如 when 中的-en。如果采取这种宽式对比态度,上面的对比程序就需要加以修正。我们现在从“不定位、定位、双项”等概念出发,再考虑到可对比程度问题,可以发现美国结构语言学后期的学者实际上采取的就是宽式对比态度。Harris(1951, appendix to 16.31)不仅从 why、when……等的对比中切分出 wh-、-y、-en……语素,还进一步从 was 和 were 的对比中切分出 w-(表过去时)、-as、-ere 这样一些语素。Wells(1947)用同样的方法从 him、them、whom 的对比中切分出 hi-、the-、who-和-m 这样一些语素,其中的-m 表示宾格或受格 (accusative)。

从我们前面给出的 where 的对比矩阵看,如果对 where、what、why、when、there、here 不作切分,那么是 6 个语素,进一步切分后则是 7 个语素,而且进行切分以后,像 th-、h-等,还必须在语素库中说明是剩余语素,只能和-ere 结合,这是不利的方面。有利的方面在于切分后可以看到 wh-的聚合关系和-ere 的聚合关系。

至于能否把对比放宽到第 4 个条件,也可以进一步研究,或许从词源研究的角度看有一定的道理。

从以上分析还可以看出,满足双项对比的不一定满足不定位对比,比如 where。那么不定位语素是否在不同的位置上总可以找到满足双项对比的实例呢?关于这个问题我们还没有进行系统调查,如果是这样,以上关于对比程序和对比性质的表述就可以进一步简化。

思考问题

1. 线性原则有哪几种理解方式?不同的理解方式有什么区别?
2. 有没有一个普遍的发现程序来识别所有语言的语素?

2. 线性序列的生成程序

在对比原则清楚以后,我们再来看线条性的实质。

前面我们说 Saussure 的线条性的第二种含义是符号线性原则。符号的连接是否完全是线性的? 如果不是,有没有线性因素? 如果有,线性的本质是什么? 这些问题还没有完全弄清楚。但有一点是不可否认的,在理解句子的过程中,能指是按照时间先后一个接一个输入的,而在生成句子的过程中,能指是按照时间的先后一个接一个说出的,这个性质对语言的理解和生成肯定起到关键作用。因此,认识线条性的本质,是语言研究中最重要任务之一。

和早期结构主义相比较,从 Bloch、Harris 开始的后结构主义一个重要的特点就是展开发现程序和形式化描写工作,其中 Harris 是最主要的代表。替换作为一种分析程序,在后结构语言学中占有重要地位。下面我们先分析替换的性质。

2.1 替换与连接

前面说过,Saussure 给出了提取线性序列中单位的基本方法:对比。对比从另一角度看也就是替换。Saussure 没有或者说很少直接涉及生成概念,如果我们注意观察儿童学习语言的过程,可以看出替换实际上也是生成新的组合的手段。以汉语为例:

喝牛奶

拿果汁

喝可乐

拿饮料

儿童能够很容易从上面的组合中替换出句法单位或语符,生成下面新的组合:

喝果汁	拿牛奶
喝饮料	拿可乐

Harris 的 *From morpheme to utterance* (1946) 的主要思路就是在论证可以通过替换生成新的话语。我们认为替换或对比无论在提取语符还是生成话语方面都是必不可少的,不过,替换不是生成上面新的组合的唯一手段。除了说“喝果汁”的生成过程是用“果汁”替换“喝牛奶”中的“牛奶”,我们也可以说“喝果汁”的生成过程是“喝”和“果汁”进行初始连接。

我们把由两个语符构成的片断称为组合基式,简称基式。基式是由两种基本方式生成的:替换和连接。替换和连接是有联系的,语符的连接和语符的替换是同一个问题的两个方面。当我们说“喝、拿”可以连接在“饮料”前,也就断定了“喝、拿”在“饮料”前面可以相互替换。表面看起来,这是从不同的角度看问题的两个概念,但在后面的讨论中我们可以看出,当组合是由两个以上的语符构成时,和连接、替换相关的生成过程会有不同的结果,会导致我们对线性结构的不同认识,并进一步会导致我们对深层结构初始顺序的不同理解。我们认为,争论什么是深层结构序列的一个重要原因就是没有区分连接和替换。

2.2 扩充与扩展

语符的替换和连接已经可以形成新的组合,但这种两个语符的组合的结果是有限的。拿替换来说,如果一个语言有 n 个语符,通过替换形成的新的两个语符的组合按照数学中的排列组合规则不会超过 n 的 2 排列。要从有限的语符生成无限的组合,必

须要用较长的组合去替换语符或用较长的组合替换较短的组合,这就是扩展。实际上 Bloomfield(1933)直接成分的概念,就是这种思路。Harris(1946)用替换关系把聚合关系程序化,并用替换关系进一步解释组合关系。当时 Harris 没有正面讨论组合关系和替换关系哪一个更初始,他认为通过语素的延长替换,即用更长的语素序列替换较短的语素序列,就可以生成无限的话语。比如:

写/书→读过/侦探小说

以上“读过”是“写”的扩展,“侦探小说”是“书”的扩展。扩展是句子生成的一种主要方式,不过在具体操作中,扩展前和扩展后是否需要在语类上相同,如果要求相同,同一性的标准怎样确定,会遇到一定的困难。严格地说,两个语素序列的语类要相同,应该是所对应的语素类都相同,这一点 Wells 已经提到过。Wells 指出(1947,3):

假定有一个序列 S,这 S 序列所属的序列类被确定为:凡属这个序列类的序列,它们的第一个语素都跟 S 序列的第一个语素一样属于相同的语素类,它们的第二个语素都跟 S 序列的第二个语素一样属于相同的语素类,以此类推;所以某个序列类中所有的成员都含有相同数目的语素。

但从 Wells 对扩展的定义和用例看,Wells 并没有把扩展限制在严格的同类上,因为模型和扩展式的语素数目可以不同,因此也不可能确定对应语素,更不可能限制对应语素的语类相同。Wells 关于模型(model)和扩展(expansion)的概念是这样解释的:

有时两个序列能出现在相同的环境里,即使它们的内部结构不一样。当这两个序列中有一序列比另一个序列长或者跟它一样(包含的语素更多或者相等)而结构不同(不属于完全相同的序列类)的时候,我们就把这个序列叫做另一序列的扩展,而把另一个序列叫做模型。假如 A 是 B 的扩展,那么 B 就是 A 的模型。(1947,4)

凡是在两个语素序列中,有一个序列是另一个序列的扩展,那么,它们的模式就是相同的。(Well,1947,1)

Wells 这里所说的“相同的环境,即使它们的内部结构不一样”,我们认为实际上就是模型和扩展式分布相同,内部结构可以不同。在具体操作中,Wells 的“相同的环境”会遇到标准宽严的问题。让我们来分析 Wells 举的英语实例 The king of England opened Parliament(英国国王召开会议):

1. **Model** John worked
2. **Expansion** The king of England opened Parliament

The king of England 可以看成是 John 的扩展,这里的 John 不是 The king of England 的成分(被扩展的项目也可以是包含在扩展项里的,这就成了陆志韦讨论的用于提取词的扩展法),但功能相同,或者说 John 和 The king of England 的分布相同。同样,opened Parliament 可以看成是 worked 的扩展,这里的 worked 不是 opened Parliament 的成分,但功能相同,或者说 worked 和 opened Parliament 的分布相同。Wells 提到这种分布是指总体分布(Wells,1947,5),英文原文是这样的:

It helps the IC-analysis to show that the sequence be-

ing analyzed is an expansion, but only if it is an expansion of the same shorter sequence in all, or a large proportion, of the environments where the shorter sequence occurs.

这里要求模型和扩展在总体分布上相同,其实就是语类相同。由于扩展可以是内部结构不一致的,成分也可以不一致,如何判定模型和扩展式在总体分布上相同,仍然会有困难。如果我们把从扩展式到模型的替换过程称为还原替换(简称还原),这种困难就很容易看出来。比如对于下面的组合:

三所/学校

我们很难从这三个语符的组合还原出两个语符的模型,这需要判定“三所”和哪个语符同类,更一般地说,需要判定数量组合和哪种语类同类。通常认为数量组合和名词同类,但两者分布上并不完全相同。再考虑下面的组合:

从北京/来

这一组合的模型也很难还原,这需要判定“从北京”和哪个单位同类,更一般地说,需要判定介宾短语和哪个语类同类。通常可能认为介宾短语和副词同类,但分布上也不完全相同。

扩展可能确实反映了人类语言类推的一种能力,是很重要的概念,但如何解决扩展所涉及的语类同一性问题,还需要做更进一步的分析。下面我们把扩展的概念加以限制,并在此基础上定义一种扩充替换(简称扩充)的概念,然后从替换和扩充的关系来认识扩展。我们所谓的扩充是把延长替换限制在同中心上,即模型和扩充式的核心不变。比如:

被扩充的单位	写	书
用于扩充的片断	写过	英语书
扩充后的结果	写过英语书	

“写过”是“写”的扩充，“写过”的组合选择和“写”相同，仍然是“书”；“英语书”是“书”的扩充，“英语书”的组合选择和“书”相同，仍然是“写”。通过这样的限制，扩充就可以直接从具体语符的组合开始，不一定要回到模型和扩展式是否分布相同的问题上来。

如果我们把同长度替换和扩充的作用结合起来，就能解释 Wells 的部分扩展。比如：

- 小马 (连接)
- 小狗 (替换)
- 大狗 (替换)
- 白狗 (替换)
- 小白狗 (扩充,对“小狗”的扩充)

从“小马”到“小狗”有替换关系，从“小狗”到“大狗、白狗”有替换关系。从“小狗”到“小白狗”，是扩充，即在“小狗”中用更长的片断“白狗”替换“狗”，“白狗”和“狗”是同中心的。从“小马”到“小白狗”，是扩展。扩展实际上包含了替换和扩充两个程序：

- 小马→小狗 (替换)
- 小狗→小白狗 (扩充)
- 小马→小白狗 (扩展)

由于替换关系和扩充关系容易得到判定，可以把替换和扩充看成是初始概念，扩展可以通过替换和扩充定义。当然我们也可

以用扩展作初始概念,说替换是同长度扩展,扩充是同核心扩展。这样看来,替换、扩充和扩展在初始性上似乎没有区别,但我们前面已经谈到,扩展和模型的同—性不容易判定,所以在讨论关键的理论问题时,我们仍然以替换和扩充为初始概念。在很多情况下,扩展仍然是一个重要概念。

与扩充相反的逆过程是减缩。前面已经谈到还原,减缩是中心相同的情况下的还原。比如:

扩充式	写过英语书	
扩充式的两个片断	写过	英语书
两个减缩片断	写	书
减缩结果	写书	

减缩在后面的讨论中是一个很有用的概念。下面把扩充与减缩、扩展与还原作一个比较:

扩充式	写过/英语书
减缩式	写/书

扩展式	写过/英语书
还原式	看/报

扩充和扩展可以生成无限组合,但扩充和扩展的生成能力并不是充分的,即扩充和扩展不能生成所有的组合。比如在上面的扩充过程中,“写书、英语书、写过”等两个语符的组合并不是扩充生成的。一般地说,用两个或两个以上的语符组成的序列去扩充或扩展一个语符,就具有句子的生成性质,但是,两个语符这一序列本身却不是扩充或扩展能够生成的,而是前面所说的初始连接(当然要以特定的聚合关系为条件)。因此,扩充和扩展并不是生成句子的充分程序。在自然语言中,首先要有

一些最基本的组合,比如:

张三睡觉	(传统语法的主谓词组)
买酒	(传统语法的述宾词组)
好酒	(传统语法的偏正词组)
苹果李子	(传统语法的并列词组)
请他来	(传统语法的兼语式)

对这些初始组合中的语符进行扩充,才可能生成更多的句子。这些初始组合是语符的直接连接形成的,于是我们说,语符的连接是生成句子的另一种程序。

2.3 叠加

替换和初始连接一般都是在两个语符之间进行的,少数是三个单位。至于更长的组合,前面讨论的一个办法是在两个语符组合模式的基础上通过扩充生成:

写小说

→写[武侠小说] (用“武侠小说”扩充“小说”)

→[写过][武侠小说] (用“写过”扩充“写”)

正因为有这种扩充的历史,因此这样的组合可以通过扩充的逆过程减缩还原成两个语符的组合:

[写过][武侠小说]

→[写过][小说] (用“小说”减缩替换“武侠小说”)

→[写][小说] (用“写”减缩替换“写过”)

但下面的组合不可能通过减缩替换还原为两个语符的连接：

[一匹][马]
→* [一][马]
→* [匹][马]

以上“一”或“匹”都不能从前面直接和“马”连接。“一马”在一定的条件可以成立，分布是有条件的，是另一个问题，这里暂时不讨论。如果我们尝试用 Wells 的扩展来解释问题，仍然会遇到麻烦。考虑下面的平行实例：

黑马
西藏马
杂交马
一匹马

前三例都是两个语符单位连接形成的组合基式或初始组合，如果认为“一匹马”是这些初始组合的扩展，就需要回答“一匹”和“黑、西藏、杂交”等分布是否相同。比如，如果我们说“一匹马”是“西藏马”的扩展，就会进一步追问“一匹”和“西藏”是否分布相同的问题。前面已经提到，判定扩展式和原型是否语类相同是一个相当复杂的问题。

为了绕开这个麻烦，让我们还是直接回到扩充所遇到的问题上来。扩充不能生成“一匹马”这样的组合，多次连接却能够。“一”和“匹”通过连接组合成“一匹”，“一匹”和“马”再通过连接组合成“一匹马”：

一 + 匹 → 一匹
一匹 + 马 → 一匹马

我们把这里的第二次连接称为叠加。叠加是基于连接的一种手段,扩充和扩展是基于替换的一种手段。和连接的不同之处在于叠加已经涉及层次,“一匹马”的叠加层次是“[一匹]马”,而不是“一[匹马]”。可以看出,凡是三个或三个以上的语符都不可能单线性连接,必须考虑层次,这是我们把连接和叠加分开的重要理由。再考虑下面的实例(“桌”在口语中有儿化):

一张圆桌

在该组合中,除了“一张”不是扩充来的,“圆桌”也不是扩充来的,因为“一张桌”是不成立的。这里的生成过程都是使用连接和叠加,不涉及替换:

一	+	张	→	一张	(连接)
圆	+	桌	→	圆桌	(连接)
一张	+	圆桌	→	一张圆桌	(叠加)

“一张”是叠加在“圆桌”上的,不是叠加在“圆”上的。下面这些组合也只能是连接和叠加生成的:

从上海来
坐在沙发上

因为这两个片断都不能缩减成两个语符的组合。

由此可见,叠加是比扩充生成能力更强的线性生成程序,叠加不仅能在基式的基础上生成扩充所能生成的组合,而且能够在基式上生成扩充不能生成的组合。

基式一般都是由两个语符构成的,还有个别情况,也可以看成是在三个语符之间进行初始连接,如:

称他小王

这类实例的生成过程解释成在两个语符的基式上再用第三个单位叠加也勉强可以,但和“一匹马”这样的组合是不同的。“一匹马”是“一”和“匹”先初始连接,然后“一匹”和“马”再叠加,其中“一匹”作为初始连接的组合,是合格的组合,但“称他”如果要算成初始连接的组合就很困难,因为“称他”不是独立片断或话语片断。“他小王”也不是合格的话语片断。这类情况也不是在两个语符的基式上进行扩充形成的。还有一些实例,解释成初始连接后再叠加也存在困难:

请他来

叫他小王

骂他混蛋

送他书

除了这些特殊情况,一般情况下三个和三个以上的语符的组合都可以看成在初始连接的基础上再进行扩充或叠加而形成的。

在基式的基础上通过等长替换(同长度替换)可以生成更多的基式。在基式的基础上通过扩充和叠加可以生成无限的组合。

前面说凡是扩充能生成的组合,叠加也能生成,那么扩充可不可以用叠加解释或定义?考虑下面言语片断生成过程:

小桌→小[方桌]

这是用“方桌”替换“桌”,生成更长的组合。这样的扩充程序完成的结果表面上似乎可以通过叠加来完成:

小+方桌→小[方桌]

即表面上是在“方桌”前面叠加“小”来形成同样的结果,但是,为什么“小”可以用同样方式加在“方桌”和“桌”的前面,有一个可替换性或聚合性的问题,即必须知道“桌”和“方桌”在“小”的后面是可以替换的,“桌”和“方桌”是一个聚合类。对比下面两个叠加,一个成立,一个不成立:

桌子→方桌子→小方桌子

桌子→抬桌子→*小抬桌子

显然,“抬桌子”和“方桌子”不是一个聚合类,不能相互替换,所以“小”能加在“方桌子”前,但不能加在“抬桌子”前。可见,是否可以叠加又是以聚合类为条件的。再考虑:

卖[白薯]

卖[烤白薯]

*卖[吃白薯]

在“卖”的后面,“烤白薯”和“白薯”有可替换性,“吃白薯”和“白薯”没有可替换性。

再考虑一个具有递归性的例子:

拜访[老师]

拜访[专家]

拜访[病人]

拜访[懂数学的老师]

拜访[在北京大学教书的老师]

拜访[懂数学的]

拜访[懂数学、有名气的老师]

拜访[懂数学、有名气的]

.....

拜访[NP]

以上第 2 个直接成分如果是词,可以通过记忆来建立连接规则,使得能够连接的都是听说过的,但第 2 个直接成分是可以扩展的,因此可以是无限长,不可能通过记忆来完成。在哪些条件下“拜访”可以叠加上去,需要知道替换条件,即上面方括号中的成分都是相互可替换的,是我们通常所说的 NP。因此,语符和语符的连接尽管可以通过记忆来实现,但语符和组合以及组合和组合的叠加则必须依赖可替换性。可以说,除了连接和替换,扩充和叠加也是生成句子的两个独立的基本程序,但这两个程序都具有可递归性。这两个程序是相互依赖的,涉及三个单位和三个以上单位的组合时,就会用到两个程序,这两个程序并不能相互推导出来。这两个原则相互依赖但又不能相互推导,这是线条性的一个重要特点。如果不考虑替换条件,这两个程序可以统一起来。

当语符首次组合的时候,使用的是连接,前面把首次连接形成的组合称为基式。现在我们可以进一步把基式定义为:

能够独立出现的最短组合。

下面这些组合不能再短,都应该看成基式:

大房子
卖房子
房子大
跑得快
称他小邵

下面的组合尽管是最短的,但不成立,因此不是基式:

- * 房子卖
- * 跑得
- * 称他

下面箭头左边的组合不能用减缩替换(扩充的逆过程)还原出基式,但能递减(叠加的逆过程)还原出基式,所以不是最短组合,因此也不是基式:

- 房子卖了→卖了
- 一匹马→一匹
- 买圆桌→圆桌
- 写在纸上→在纸上→纸上

我们认为基式中成分的顺序是语符线性组合的基本顺序。

当基式延长时,前面提到可以用叠加和扩充。对基式中的单位进行扩充,自然就形成了直接成分,形成了层次。叠加也会形成层次。叠加有两种方式,一种是加在整个组合上,一种是加在某个单位上。以“语法书”的左线性叠加为例:

- 语法书
- [生成语法]书 (加在成分“语法”前)
- 线装[语法书] (加在组合“语法书”前)

这种叠加和扩充是等价的,比较:

线性叠加		扩充	
叠加前	叠加后	扩充前	扩充后
生成语法	[生成语法]书	语法书	[生成语法]书
语法书	线装[语法书]	线装书	线装[语法书]

对于长于基式的组合的扩充,原则是相同的。一般地说,在基式上通过扩充替换得到的组合,也能够通过叠加得到。但反过来不成立,因为扩充总可以还原到基式,但叠加不一定能够还原出基式。比如:

[《阿 Q 正传》][我已经读过了]

最终还原不出“《阿 Q 正传》读”这样的基式,“读《阿 Q 正传》”才是基式。

可见有层次的线性结构有两种,都可以看成是直接成分的组合,都可以进行层次分析。下面相同层次的直接成分用相同的括号,左括号的右下方表示括号中短语或词的语类:

$$\begin{aligned} &\langle_N \text{我}\rangle \langle_{VP} (\text{adv 已经}) (\text{VP} [\text{VP 读过了}] [\text{N} \langle \text{《阿 Q 正传》} \rangle]) \rangle \\ &\langle_N \langle \text{《阿 Q 正传》} \rangle \rangle \langle_S (\text{N} \langle \text{我} \rangle) (\text{VP} [\text{adv 已经}] [\text{VP 读过了}]) \rangle \end{aligned}$$

但这两种线性结构的层次的生成过程并不相同,一种是在基式的基础上可以通过扩充得到的,因此这样的扩充式也可以还原出基式,一种是在基式的基础上必须通过叠加得到的,这样的叠加式不能还原出基式。也就是说,通过叠加得到的线性序列,有一部分是可以通过扩充得到,有一部分不能通过扩充替换得到。我们后面会讨论,能够通过扩充得到的基式,基本上等价于 Chomsky(1965)的深层结构,不能通过扩充得到的线性序列,在 Chomsky 的生成语法中是放在转换规则中处理的。其实这一部分由叠加而不是由扩充形成的线性序列也可以通过直接成分的规则来处理。从这里我们发现,Chomsky 的短语结构规则并不是直接成分的形式化,而是扩充的形式化。

尽管扩充生成的组合用叠加也能生成,但扩充仍然是独立的生成程序,表现在两个方面,一方面,基式的扩充规定了可减

缩线性结构的集合,因此可作为确定某一类深层结构的标准。另一方面,扩充并不改变由基式规定的模型,而在保持模型不变的情况下进行扩充是句子生成过程中的一个基本要求。在生成很长的组合的情况下,有时候直接用扩充来表述生成过程更容易显示组合的线性属性。比如,如果已经有了这样的组合:

邵永海送给〈学生〉一句话

我们可以通过对“学生”进行扩充生成以下组合:

邵永海送给〈(大一学生)〉一句话

邵永海送给〈(大一[优秀学生])〉一句话

如果要用叠加来表述后一个实例的生成过程,需要把句子拆分成语符后重新组合,然后一个语符一个语符依次叠加:

→[优秀学生]

→(大一[优秀学生])

→〈送给(大一[优秀学生])〉

→{〈送给(大一[优秀学生])〉一句话}

→邵永海{〈送给(大一[优秀学生])〉一句话}

随着要生成的句子的拉长,直接叠加的表述方式不容易看出基式和叠加式的同模关系,当然这仍然不能否认叠加是比扩充更强的线性生成手段。

现在我们可以说,在线性生成过程中存在连接、替换、扩充、叠加四种初始程序。连接与替换形成基式。在基式的基础上再使用叠加和扩充形成更复杂的结构。比较起来,叠加是比扩充更强的生成方式,只要是扩充能得到的序列,用叠加一定能够得

到,但用叠加得到的序列,扩充不一定能够得到,比如“一匹马”。用扩充得到的序列可以通过减缩替换得到原形,通过叠加得到的序列不一定能够通过减缩替换得到原形。这会决定对转换使用的态度。如果在扩充结构的基础上使用转换,转换规则就会用得很多,如果在叠加的基础上使用转换规则,转换规则就会用得很少。在后面的分析中,这四个程序是认识其他理论方法的重要出发点和基础。

Harris(1951)曾考虑只用分布来描写语法,现在来考虑分布和以上四个程序的关系。当我们说“一”和“匹”可以连接,也可以说“一”可以分布在“匹”的前面,或说“匹”可以分布在“一”的后面,可见连接可以用分布来解释。进一步,“马”可以和“一匹”叠加,也可以说“马”可以分布在“一匹”后面,所以叠加也可能用分布来解释。

有些“替换”用分布来解释也可以。以“买一瓶酒”为例,该组合的各个位置可以进一步替换,生成更多的组合。比如:

买一瓶酒	买一瓶酒	买一瓶酒	买一瓶酒
喝一瓶酒	买两瓶酒	买一碗酒	买一瓶水
送一瓶酒	买三瓶酒	买一壶酒	买一瓶奶

在“买一瓶酒”这样的线性序列中,特定位置的 4 个语符都可以进一步替换。比如在第一个位置,“买、喝、送”相互可以替换,也就是说在这一线性序列中“买、喝、送”有相同的分布,即“买、喝、送”都可以分布在“酒”前面,所以以上线性替换从根本上说又可以概括成分布。更进一步,以上的“酒”也可以用“好酒”扩充,我们也可以说“好酒”可以分布在“一瓶”后面,所以扩充也可以用分布解释。如果人类的语言都符合这样一种线条性原则,最终都可以用分布解释,那么人类语言的生成机制就可以统一到一个相当简明、直观的程序。我们认为线性原则是后来发现程序试图用分布来控制解释组合规则的一个重要理论背景。不过,

前面讨论过的有些线性组合是不能通过减缩替换还原的,所以分布还不能解释扩充和叠加的区别。分布是有用的概念,但后面讨论其他问题时可以看出,如果用分布统一连接、替换、扩充、叠加等概念,会掩盖线性生成的本质。

以上讨论的连接、替换、扩充和叠加,所生成的组合都是直接成分的线性组合。如果在话语的生成过程中坚持只用连接、替换、扩充和叠加,则生成的话语序列都满足有层次的线条性。如果我们能够证明只用连接、替换、扩充和叠加不能生成语言中的有些句子,就可以证明语言的句子不完全是线性的。后面分析转换规则的必要性时我们将讨论线性原则的不充分性。

思考问题

1. 生成线性序列至少需要哪些程序?
2. 线性组合的基式是否可以全部归纳为两个语符的组合?为什么?

3. 单位和规则

以上在讨论线性序列的生成手段替换、连接、扩充、扩展、叠加时,是以已经获得句法单位为前提的。实际上获取句法单位是相当复杂的过程,这方面存在的困难也是结构语言学、转换生成语法、功能语法等所遇到的一些难点问题的起因。如果两个语素的组合是规则组合,就是两个句法单位,如果是不规则组合,就是一个句法单位,可见单位和规则的判定是同一个问题。

3.1 词作为句法单位遇到的问题

Bloomfield(1933)把词定义为最小的自由形式,实际上这个定义只反映词作为生成单位的部分性质(陈保亚,1992)。徐通锵(1994)开始考虑把汉语的句法单位建立在字的基础上,这方面讨论还在进行中,但词的问题已经逐渐显现出来。

如果我们坚持句法单位的有限性原则,则 Bloomfield 的词的定义已经不符合有限性原则。以英语为例:

seventh
twenty seventh
thirty seventh
forty seventh
.....
one hundred and twenty seventh
two hundred and twenty seventh
.....

这里的派生模式可以概括为：

X-th

显然, X-th 中的 X 是无限的。

Bloomfield 曾经提出过一个“长单词”的概念(1933, 11. 5, p219):

我们写出的 the boy's (这男孩子的), 好像有两三个词, 但是, 严格来讲, 那只是一个词, 因为直接成分是 the boy 和属格 [-z], 后者是一个黏附形式; 这种情况在 the king of England's (英国国王的) 或 the man I saw yesterday's (我昨天看见的那个人的) 一类的例子中显得更清楚, 其意义表明 [-z] 是跟它前面的整个短语结合在一起的, 因此这两部分就联成为一个长的单词。

如果承认 Bloomfield“长单词”的观点, 也就等于承认词可以是无限的。这也不符合有限单位原则。从规则的性质看, 其实英语中的 X-th 和 X's 的实质是一样的。

Wells 确定了一条原则, 每一个词都应该是一个成分(1947, 40)。这个定义还不是很明确。根据 Wells 这个原则, 确定词在英语中有时会遇到 Bloomfield 所遇到的困难。拿 Wells 的例子来说:

the president of the bank's daughter

bank's 在这里并不是一个成分, 因为这里的第一个层次是(符号“/”表示层次的界限):

the president of the bank'/s daughter

通常我们说英语的所有格“s”是词尾,由于这里的“s”是直接控制 the president of the bank 的,就必须把 the president of the bank' 看成是词。这是一个非常奇怪的现象。这也是 Bloomfield 关于词的定义所遇到的最为困难的地方之一。

Wells 还提出了另一种选择的可能性,即把's看成是词。这样一来,就可以满足每一个词都是一个成分的原则(1947,40)。但按照 Wells 的这个选择,词就不一定有重音,同时也不满足自由(单说)这一条件。Wells 的这种确定词的办法实际上是一种扩展法。

3.2 单位及组合的可推导性

在汉语中,词也遇到很多问题(陈保亚 1999,6.2)。考虑下面三组实例:

第一、第二、第三、第四、第五、第六

老张、老王、老李、老陈、老杜、老刘……

桌子、椅子、房子、裤子、凳子、袖子……

从词法研究的角度看,上面三组实例都是派生构词,差别在于第一、二组是带前缀的,第三组是带后缀的。我们要追问的问题是,哪一组实例更容易掌握?显然是第一、二组,因为第一、二组的派生方式是可以周遍类推的,不需要死记。而第三组的派生方式是不可周遍类推的。这种区别涉及类推的本质,是确定单位和规则必须要考虑的问题。有一点是很清楚的,如果“第 X”是词,就必须假定词是无限的。

前面说 Saussure(1916)最先提出提取单位的基本原则是对

比,后来 Bloomfield(1926,1933)以及结构语言学家深入讨论了对比。如果给出一定的限制,用对比原则提取语素基本上是可行的(陈保亚,1997,3)。语素是话语分析中非常重要的一种单位,问题是语素可不可以作为生成话语的基本单位? Harris(1946)曾经设计了一种话语生成程序,就是要直接从语素通过替换生成话语。这一话语生成程序包含了一个前提:生成话语的基本单位是语素。

如果 Harris 的话语生成程序能够成功,就可以绕开上面以词为单位生成话语所遇到的问题。下面我们将论证这个程序的生成能力过强,同时将论证以 Harris 的思想为基础的短语结构语法的生成能力也过强。这种过强的生成能力从根本上说不是因为没考虑语义,因为即使考虑意义,仍然存在生成能力过强的问题。其实 Harris 并没有说他的模型最终不考虑语义,而是想在语义条件相同的情况下考虑语法的性质。Harris 模型存在的根本问题是:如果把语素作为句法单位,无法区分规则组合和不规则组合。

Harris 的基本思路是:语素组合的类可以从语素的类推导出来。这是自 Bloomfield(1933)以来关于单位和短语的推导关系的一次大的转变。Bloomfield 的语类推导思想是建立在词类基础上的:

一个短语的形类通常是跟包含在短语中的某个词的形类相同。例外的情况就是离心结构的短语,而我们已经看到这也能根据词类来下定义。因此,短语的句法形类能够由词的句法形类而得来:也就是句法的形类最容易根据词类来描写。(Bloomfield,1933,12.11,p241)

Harris(1946)是要把句法单位建立在语素上,这样能够绕开识别词所遇到的困难。为了分析 Harris 分析方案的可行性,我们

可以拿汉语的实例来表达 Harris 的思路：

A	V	VP
近	看	近看
细	看	细看
小	看	小看
轻	拍	轻拍
重	拍	重拍

Harris 的意思是说，A 类语素加 V 类语素是 VP 类语素组。我们认为这种语类推导思想是后来转换生成语法的短语结构规则和语杠理论 (Chomsky, 1957, 1970, 1986) 暗含的一种基本思想：即只要直接成分的语类确定了，短语的语类就可以推导出来；只要短语的语类确定了，直接成分的语类也可以推导出来。

Harris 的语类推导思想是想通过分布解释话语生成的递归性，但把语类推导建立在语素这一级单位上，却会遇到困难。首先遇到的问题是不能有效区分规则组合和不规则组合。语素在线性方向的组合有两种根本不同的行为：有规则的活动方式和无规则的活动方式。比如“白车”是指白颜色的车，其整个片断的意义可以通过“白”和“车”的意义加上组合关系推导出来，这是语素有规则的组合。但“白菜”却不是指白颜色的菜，整个片断的意义不可以通过“白”和“菜”的意义再加上组合关系推导出来，这是语素无规则的组合。很明显，不规则语素组是需要记忆的，也不满足语类推导，这类语素组在数量上应该是有限的。有效区分规则语素组和不规则语素组是语言习得和自然语言理解的基础，也只有规则语素组才可以进行语类推导。如果把语素作为生成话语的基本单位，不能说明为什么“白车”的意义可以推导，而“白菜”的意义不能推导。

把语类推导建立在语素这一级单位上，遇到的另一个问题

是类推的不一致问题。以汉语语素“视”为例,我们遇到的第一个问题是怎样确定“视”的语类,因为这不是一个自由形式。即使我们暂定“视”的类是 V,由这样的语素组成的语素序列的类是不可以推导的,更准确地说是类推方式不一致。比如:

	左项语素的类	右项语素的类	语素组的类
远视	A	V	N
近视	A	V	N,A
斜视	A	V	N,V
重视	A	V	V,A

以上前后两个语素的语类相同,但语素组的类并不一致。即使我们把“视”看成某个新的类,或者先不给它定语类,也仍然面临上面语素组类推不一致的情况。可见,如果把语素作为生成话语的单位,语类推导就不能实现;如果没有语类推导,生成话语的规则就很难确定,生成语法的短语结构规则和语杠理论也就失去了基础。

把语素作为生成话语的基本单位,还会造成错误类推,比较下面的语素组合:

本儿、花儿、凳儿、瓶儿
* 笔儿、* 筷儿、* 纸儿

原因在于“X 儿”并不是可以周遍类推的组合,如果把“X 儿”看成是句法单位,就不会出现过分的类推。

无论是语义推导还是语类推导,都是从单位到组合的一种推导过程,我们认为可推导性是自然语言生成性的核心,说话者和听话者双方在享有共同的单位和规则的前提下,说话者之所以能够说出从来没有说过的组合,而听话者能听懂从来没有听说过的组合,都依赖了可推导性。生成话语的单位首先应该满足可推导性,从上面的分析看,这种可推导性所依赖的单位不是

语素。概括地说,把语素作为句法单位存在两方面的问题:

1. 不能区分可类推和不可类推。
2. 会造成过分类推。

生成语法没有正面讨论单位和规则的关系,一些涉及单位和规则的讨论主要出现在 Chomsky(1981)支配理论以前,多出现在 Chomsky(1957,1964,1965)、Chomsky 和 Halle(1968)的论著以及生成语法中关于词汇派 (lexicalist position) 和转换派 (transformational position) 的争论中。一般认为直接成分理论是 Chomsky(1957)短语结构规则的基础。更准确地看,我们认为 Harris(1946)的从语素类推话语的模型是 Chomsky(1957)的转换生成语法中的短语结构规则的更直接更主要的思想基础,Chomsky 更进一步把语素话语模型形式化了,并把语类替换改变了一个方向,从抽象的句子开始不断改写直接成分,最后改写到语素。Chomsky (1957)的语素(morpheme)和词(word)之间有一种关系,这种关系体现在转换生成语法所包含的三个部分之间(Chomsky,1957,p46):

Phrase structure

Transformational structure

Morphophonemics

Chomsky(1957,p46) 认为这三个部分的关系是:

To produce a sentence from such a grammar we construct an extended derivation beginning with *Sentence*. Running through the rules of F we construct a terminal string that will be a sequence of morphemes, though not

necessarily in the correct order. We then run through the sequence of transformations $T_1 \cdots T_j$, applying each obligatory one and perhaps certain optional ones. These transformations may rearrange strings or may add or delete morphemes. As a result, they yield a string of words. We then run through the morphophonemic rules, thereby converting this string of words into a string of phonemes.

可以看出,Chomsky 语素和词的关系体现在生成过程中,短语结构(Phrase structure)的改写规则生成语素序列[包括暂拟语素(tentative morpheme)],转换结构(Transformational structure)生成词序列。在后面的讨论中我们会谈到,短语结构的生成能力和扩展替换的生成能力是一致的,要准确描写一个语言的短语规则或扩展规则,必须要区分规则组合和不规则组合。

Chomsky(1965)基本没有用 morpheme 这个单位,而是用 formative 作句法单位。其实 formative 基本上和 morpheme 是相同的。先让我们来证明这一点。Bloomfield(1926, 13)曾经提出并定义过 formative:

A bound form which is part of a word is a formative.

Bloomfield 的 formative 只是词的有意义的构成部分,是黏着形式。Chomsky(1965)、Chomsky 和 Halle(1968, p8)的 formatives 含义不限于黏着形式,而是包括了暂拟语素(tentative morpheme)、黏着语素和自由语素,这和 Chomsky(1957)的语素范围基本是一致的。从 Chomsky 的用例看,formative 也相当于语素。短语结构规则生成 phrase maker,所插入的 formative 都相当于语素。比如 Chomsky 和 Halle(1968, p8)的实例:

(3) We established telegraphic communication

其 formatives 是:

+ we + + establish + + past + + tele + + graph + + ic + +
communicate + + ion +

两个十号之间的都是 formative。有几点值得注意:

1. formative 可以是暂拟的, 如 past。

2. Chomsky(1964, p65; 1965, p65) 把 formative 分成词汇项目 (lexical item) 和语法项目 (grammatical item)。比如在 sincerity may frighten the boy 中, sincerity、boy 是词汇项目, the 是语法项目。以上的例子中, we、establish 等都是词汇项目, past 是语法项目。语法项目还包括词类标记 (class markers) 以及语素间隔标记、词界 (word boundary) 等音渡 (juncture) (Chomsky 1964, p66)。

3. Chomsky(1957) 的 morphemes 序列是短语结构生成的, 经过转换变成了词。Chomsky(1965) 是 formative 序列构成了表层结构, Chomsky and Halle(1968, p12) 提到生成表层结构的规则要提供词的信息, 才能使用音系规则。至于怎样在规则中提供词的信息, Chomsky 没有讨论。

显然, formative 和 Chomsky(1957) 的语素是相当一致的。

Chomsky 并没有给 formative 明确的定义, 从 Chomsky 的行文中可以看出 formative 作为句法单位的重要性 (Chomsky, 1965, p65):

A grammar that generates simple Phrase-markers

such as (3)^① may be based on a vocabulary of symbols that includes both formative(the, boy, etc.) and category symbols (S, NP, V, etc.). The formatives, furthermore, can be subdivided into lexical items (sincerity, boy) and grammatical items (Perfect, Possessive, etc.; except possibly for *the*, none of these are represented in the simplified example given).

这一段明确说明了短语结构的最小单位是 formative。

现在来考虑 Chomsky(1965)的 formative 是怎样和句法发生关系的。Chomsky(1965, 2.3.3, p87)指出:

In general, all properties of a formative that are essentially idiosyncratic will be specified in the lexicon' (formative lexical item)。

Chomsky(1965, 2.3.3, p87)认为,

the lexical entry must specify lexical features indicating the positions in which a lexical formative can be inserted (by the lexical rule) in a preterminal string.

也就是说,改写规则生成语类序列,这些语类序列通过词汇插入规则插入的是 formative,这也是把 formative 作为句法单位的重要证据。这种思想在后来 Chomsky(1981)为代表的支配理论中一直没有改变。只不过把 formative 换成了 lexical item。Chomsky(1981, 1, p5)对 rule system 的各部分的关系以及各表

① (3)指 Sincerity may frighten the boy 的短语标记(phrase marker)。

达层面的关系做了很好的解释：

The lexicon specifies the abstract morpho-phonological structure of each lexical item and its syntactic features, including its categorial features and its contextual features. The rules of the categorial component meet some variety of *X-bar* theory. System (i) and (iia) constitute the base. Base rules generate D-structures (deep structures) through insertion of lexical items into structures generated by (iia), in accordance with their feature structure. These are mapped to S-structure by the rule *Move- α* , leaving traces coindexed with their antecedents, this rule constitutes the transformational component (iib), and may also appear in the PF- and LF-components. Thus the syntax generates S-structure which are assigned PF- and LF-representations by components (iii) and (iv) of (1), respectively.

根据上面已经给出的材料, lexical item 就是 formative (不包括语法成分)。Chomsky (1965) 也提到要确定词的界限才能使用语音规则, 并没有把词作为短语结构规则中的单位。

Chomsky (1965) 的 formative 是词库中的单位, 词库中的单位也就是 formative, 仍然相当于把句法的单位建立在语素上, 因此, Chomsky (1965) 的理论模型仍然存在 Chomsky (1957) 所面临的问题, 同时也面临 Harris (1946) 所遇到的问题。

下面来分析 Chomsky 和 Halle (1968) 的实例需要用到哪些规则：

(3) We established telegraphic communication

下面两个十号之间是 formative:

+we++establish++past++tele++graph++ic++communicate++ion

其中 past 是一个暂拟语素, we、establish、communicate 是自由语素, tele、ic、ion 是黏着语素。和 Harris(1946)语素类推思想不同的是, Chomsky (1957, 1965)、Chomsky 和 Halle (1968)的语类规则只是指派到词,并没有指派到语素,比如上面实例的语类指派是(这里我们省略了短语语类的指派):

$[_N \text{we}] + [_V [_V \text{establish}] [_{\text{past}}]] + [_A [_N [_{\text{tele}}] [_{\text{stem}} \text{graph}}]]] [_{\text{ic}}] + [_N [_V \text{communicate}] \text{ion}]$

不过 Chomsky (1964, p60) 从句子到语素的改写思路是很明确的:

A generative grammar consists of a syntactic component, which generates strings of formatives and specifies their structural features and interrelations.

从前面的语类指派实例以及 Chomsky 原文所给出的树形图容易看出,要生成上面的句子必须用以下改写规则

$V \rightarrow V + \text{past}$	$([_V [_V \text{establish}] [_{\text{past}}]])$
$A \rightarrow N + \text{ic}$	$([_A [_N [_{\text{tele}}] [_{\text{stem}} \text{graph}}]]] [_{\text{ic}}])$
$N \rightarrow V + \text{ion}$	$([_N [_V \text{communicate}] \text{ion}])$

实际上第一条规则和第二、三条规则有完全不同的性质,如果都作为相同性质的规则来使用,就会出现过分类推或生成力太强。比如第二条规则可以同时生成以下两个带 ic 的实例:

telegraph telegraphic
blackboard ' blackboardic

真实语言中 blackboardic 是不存在的。Chomsky 和 Halle (1968,p12)提到要通过分析把插入 formative 后的序列分析成词,再使用音系规则。我们认为要解决以上规则生成力过强的问题,在插入语素以后来解决是不可能的,必须先把 telegraphic 这样一些片断作为单位来处理,存放在词库中,因为 N+ic 构成形容词和 V+ion 构成名词是解释规则而不是生成规则(详见后面的讨论),即我们只能肯定 N 加 ic 的组合是形容词,V+ion 的组合是名词,但并不是所有的名词 N 都可以加 ic 构成形容词,也不是所有的动词都可以加 ion 构成名词。这些规则是不可以周遍类推的(参考平行周遍原则部分),这里的情况和汉语“X-儿”(花儿、心儿)有共同的地方。

以上分析是说,Chomsky(1965)、Chomsky 和 Halle(1968)的改写规则并没有区分开生成性规则组合和解释性规则组合。尽管 Chomsky 多处提到要区分规则和不规则现象,这一思想也确实能解决 Harris 的语素话语模型所遇到的类推问题,但由于没有严格区分平行规则和平行周遍规则,所以改写规则会生成上面提到的不合法片断。在传统语法中,上面第一条规则所描写的情况属于形态学的范畴,第二、三条规则所描写的情况属于构词学的范畴。后面我们会谈到,规则存在平行和平行周遍两种情况,但是传统语法没有按照规则的条件来区分这两类情况。比如 seventh 需要第一类规则来描写,但传统语法却列为构词范畴。Chomsky 也没有从规则的平行周遍性质来区分两类情况。

两类规则是认识类推性和生成性的关键,以语素为句法单位不能反映这两类规则,因此不能有效地完成短语结构的描写。比如,下面的描写数量组合的短语结构是成立的:

$QP \rightarrow Q + L$ (一匹)

这条规则对数量语素的组合是周遍的,更具体地说,数词在量词前是周遍的。下面描写汉语儿化的语素组合却是不成立的短语结构:

$F \rightarrow N + 儿$

因为什么条件下一个名词性语素可以加“儿”,条件并不清楚。如果根据这个短语结构规则生成汉语的儿化词,有很多是不成立的,会生成“·笔儿、·书儿”这样一些组合。

又比如,如果要以语素为句法单位描写前面提到过的组合:

	左项语素的类	右项语素的类	语素组的类
远视	A	V	NP
近视	A	V	NP, AP
斜视	A	V	NP, VP
重视	A	V	VP, AP

就不得不建立下面的改写规则:

$NP \rightarrow A + V$ (远视、近视)
 $VP \rightarrow A + V$ (斜视、重视)
 $AP \rightarrow A + V$ (近视、重视)

这些规则所使用的周遍条件是什么,还不清楚。如果作为生成规则来使用,带来的后果是,本来大家公认的可以周遍类推的规则 $VP \rightarrow A + V$,由于把句法单位的基础建立在语素上,现在不

能周遍类推了。而本来不能周遍类推的规则,比如 $NP \rightarrow A + V$, 如果一定要作为生成规则,则要生成出大量的错误组合。

可见,短语结构描写的语素组必须是根据可周遍类推规则组合的语素组,由于不是所有的语素组都是可周遍类推规则组合,所以在语素这个层面上建立短语结构存在问题。

Chomsky 和 Halle(1968)有一道生成手续是通过转换规则把语素序列转换成词序列,这样就可以进一步使用音系规则把词序列转换成语音序列。但是,在转换结构中,要把哪些语素组合转换成词,也需要区分可周遍类推规则组合与不可周遍类推规则组合。在英语中除了加形态的成分,可周遍类推规则组合的语素组都不应该转换成词,但是哪些语素组是可周遍类推组合,哪些不是,生成语法并没有给出一种判定方法。后面我们还会讨论转换和平行周遍的问题。

如果没有把单位和组合区别开,建立在语类推导基础上的短语结构规则的描写会造成生成力强弱悖论,具体表现为,一方面生成规则的生成能力太弱,另一方面生成规则的生成能力过强。以汉语为例,过弱的例子有:

$X \rightarrow \text{第} + \text{数字}$

通常认为“第三”这样的组合是词,“第”是词头,把“第 X”放入词库,不能反映生成能力。“第 X”是有无限生成能力的。

生成力过强的例子:

$NP \rightarrow X + \text{儿}$

$NP \rightarrow X + \text{子}$

$NP \rightarrow X + \text{头}$

这样的生成式会生成“·书儿、·笔儿、·衣子、·灯子、·钢头、

“铁头”这样一些组合。实际上哪些语素可以进入上面公式并不清楚的,需要记忆。如果要坚持语素插入,就不能避开这样的公式,必须是规则句法单位或语符插入,公式才是有效的。

如果在短语结构规则生成的语类序列中插入 formative,就会生成不合法的句子:

N [V N]

他打人

他举手

* 他随手

后来 Chomsky 开始重视规则形式和不规则形式的区分,并且系统讨论过规则和不规则问题(Chomsky, 1970)。我们认为 20 世纪 60 到 70 年代生成语法关于限制转换规则生成能力过强的讨论,本质上也是和区分规则组合与不规则组合有关系。面对 destroy 和 destruction 这两种实例的关系,当时出现了转换假说(transformalist hypothesis)和词汇学假说(lexicalist hypothesis)的争论。转换假说认为 destruction 是由 destroy 转换出来的,因此这样的派生过程是句法的过程,需要在句法部分解决。词汇假说认为 destruction 的构成是在词汇部分,应该在词汇部分处理。后来词汇假说被更多的学者接受。我们认为,转换假说的思路和 Harris 从语素到话语的思路是一致的,这种思路把从 destroy + ion 派生出 destruction 也看成是一种规则,并没有区分规则中的两种情况,即周遍的和不周遍的,即没有区分 destroying 和 destruction 在规则上的本质区别。尽管词汇假说意识到要区分这两种情况,由于没有追问单位的提取和规则的确定依据什么分析程序或原则,转换生成语法后来一直没有给出区分规则组合和不规则组合的方法,没有解释句法研究需要什么单位才能反映句法的生成性或可推导性。

尽管人们没有从正面论证为什么语素不能作为句法的基本单位,但在具体语言分析中已经碰到用语素描写句法的困难,所以目前大多数学者的做法仍然是把词作为生成话语的基本单位。不容否认,词在韵律上表现出很多规律,这方面已经有很多证据(王洪君,2000.6),但词作为生成话语的单位是否满足可推导性?这取决于我们怎样定义词,用什么办法来提取词。Bloomfield(1926,1933)对词做了深入的研究,把词定义成最小的自由形式,奠定了用单说的方式提取词的原则。这一原则能够提取单说的最小自由形式,但不能提取不单说的形式,所以在下面公认的词组中,前项都不能用单说原则提取:

再来、又来、常来……

陆志韦(1937)在单说论的基础上补充剩余法来提取“再、老、常”这类不自由形式。即只要“来”可以单说,剩下的也是词。这个办法也会遇到困难,考虑下面语素组:

将来、本来、从来……

这里的“将、本、从”并不能用剩余法提取。实际上剩余法预先假定了“再来、又来、常来”是词组,“再、又、常”是词,而“将来、本来、从来”是词,其成分“将、本、从”不再是词。也就是说,“再、又、常”并不能通过剩余法提取出来。单说论或单说论加上剩余法不能提取“再、又、常”这类不自由的词。我们也不好否定这些语素不是词,因为包含这些语素的组合是可周遍推导的,现代汉语研究中把这些语素叫“虚词”。单说论并不反映可推导性。

后来陆志韦(1957)又提出扩展法来提取词。我们认为扩展法的实质是看一个片断的直接成分是否可以被更长的同形片断替换,如果能,就是词组(陈保亚,1999,2.1.2)。但是扩展法也

不完全反映可推导性,下面的 A 式可以通过成分和组合规则推导出整体意义,有可推导性,B 不能通过成分和组合规则推导出整体意义,B 式没有可推导性,但 A 式和 B 式都可以扩展:

A	A 式的扩展	B	B 式的扩展
买书	买[一本书]	理发	理[一个发]
做饭	[做两次]饭	洗澡	[洗两次]澡

另一方面,不能扩展的片断有不少是有可周遍推导性的:

老王	小王	初一
老李	小李	初二
.....		

概括地说,有些没有可推导性的片断,可以扩展,有些有可推导性甚至可周遍推导性的片断,不可以扩展,所以扩展法不反映可推导性。扩展法反映的仍然是语素活动的自由程度。

语素活动的自由程度除了单说、扩展两个主要衡量标准外,还有转换标准。转换也不反映可推导性,下面的“买书”有可推导性,“理发、研究”没有可推导性,但都可以有不同程度的转换:

原式	第一转换式	第二转换式
买书	买不买书	书已经买了
理发	理不理发	发已经理了
研究	研不研究	

而“老张”、“第三”等有可推导性的片断,反而没有转换式。

单说论、扩展法和转换法仍然是目前区别词和词组的主要语言学方法,但这三种方法不能充分区分不规则语素组和规则语素组,不能充分反映可推导性。当然区别词和词组跟区别不规则语素组和规则语素组是两个不同的问题,目的也不一样,“自由、可单说、可扩展、可转换”也都是很有价值的概念,朱德熙

(1982)关于黏着式偏正结构和组合式偏正结构的区分,就是建立在可扩展基础上的。我们认为自由是有程度的,最自由的是单说,其次是移位(转换),再其次是扩展(延长替换),再其次是等长替换,最后是不能替换。自由的概念对于词法、句法分析相当重要。但是,自由和平行周遍是两个不完全相同的概念,平行周遍关注的是规则和不规则的制约条件。我们将在后面相关部分讨论这些问题。下面我们集中讨论我们曾提出的平行周遍原则,观察这一原则在区分规则语素组和不规则语素组的过程中会遇到什么问题。

为了后面讨论的方便,在很多情况下我们把规则语素组“白车”中的“白”和“车”称为语符^①,把不规则语素组“白菜”也称为语符。关于语符,可以给出两个定义:

最短规则语素组中的直接成分。

没有一个成分可以进行平行周遍替换的有意义的片断。

语符组合满足可周遍推导的性质。人脑或电脑中储存语符的地方称为语符库,储存规则的地方称为规则库。语符库和规则库中存放的项目数量是有限的,生成话语的基本单位和生成话语的规则是紧密联系在一起的,因此提取语符和规则时,有两个更基本的原则需要遵守:

有限原则:语符的数量是有限的,规则的数量是有限的。

最优原则:在语符的数量已经确定的情况下,规则的数量越少越好;在规则的数量已经确定的情况下,语符的数量越少越好。

^① 黄长著等把 Chomsky(1965) 的 formative 翻译为语符,本文的语符有不同的含义,详见下文。

3.3 平行周遍原则与生成规则

当我们说“白车”是规则组合,其意义可以通过语素和组合关系推导出来时,这里的“推导”、“规则”只是一种不严格的说法,并不是严格的操作程序。“近视、远视、弱视”是否可以推导出来就不容易断定,因此仅仅靠“推导”、“规则”的定义还不容易确定它们是规则语素组合还是不规则语素组合。又比如“心儿、门儿”这类儿化现象,从成分和组合关系应该可以推导出整体的意义,似乎应该是规则组合,后面会看到“X 儿”这类组合中只有理解规则,应该归入语符。

文献中对可类推或可推导性的理解和用法也是不一样的,“心儿、门儿”有的认为是可以类推的,有的认为是不可以类推的。可见类推性需要严格界定。我们需要找到一种严格的方法或程序来区分规则组合和不规则组合。前面提到的平行周遍原则,就是我们试图对规则做出判定的一种尝试。现在来具体分析平行周遍原则,观察下面语素组:

1	2	3	4
老来	老李	老虎	老板
老去	老张	老鹰	老手
老睡	老刘	老鼠	老实

第 1 列的对比关系是平行的,第 2 列的对比关系也是平行的。第 3 列的第 2 个成分是和动物有关的语素,我们说第 3 列的对比关系也是平行的。第 4 列的“板、手、实”很难找到共同特征,我们说第 4 列的对比关系是不平行的。

从以上语素组可以看出平行性包括了以下三个方面:

- 1. 被替换的部分具有平行特征。

2. 在被替换部分保持平行特征的前提下,组合关系平行。
3. 整个组合在分布上平行。

上面第 1、2、3 列满足这三个条件,第 4 列的语素组不满足这三个条件。

尽管第 1、2、3 列都满足平行对比的条件,但有很大的区别,第 1 列的平行关系可以周遍到每一个自由的行为语素,第 2 列的平行关系可以周遍到每一个单音节姓氏语素,第 3 列的对比关系却不能周遍到每一个单音节动物语素,“老鸡”就不可以说。我们说第 1、2 列对比不仅平行,而且周遍,而第 3 列对比只平行,但不周遍。不周遍的含义是,我们总可以找出平行的实例,但却不可以言说。

以上三个平行条件也可以用规则来表述,比如对于“老张、老李”这样的语素组,我们可以给出这样的平行规则:

单音节姓氏语素放在‘老’的后面,整个语素组表示尊重。

由于这条平行规则是周遍的,所以也可以说成:

所有单音节姓氏语素放在‘老’的后面,整个语素组表示尊重。

也就是说,三个平行条件等价于平行规则,平行周遍等价于含有全称量词的平行规则。为了讨论问题的方便,我们有时候也用“平行规则”指三个平行条件,用“平行周遍规则”指平行周遍条件。如果不能从某些相似的语素组中提取出一条规则,这些语素组就不算是平行的。

“老虎、老鹰、老鼠”这样一种平行关系已经显示出一定的规则性,我们可以从字面上推导出这些语素组的意义,但由于这种关系不具有周遍性,我们无法预测哪些单音节动物语素可以和

“老”组合,哪些不能,因此这类组合方式不能无限制地类推,可以说“老鼠”,但不能说“老兔”。“老来、老去、老睡”的平行关系以及“老李、老张、老刘”这样的平行关系由于是周遍的,可以无限制地类推,生成“老哭、老吃”以及“老王、老贾”等新的语素组。

从“老 X”的分析看,我们可以把语素组分成 3 种类型:

- 1. 不规则语素组:如“老板、老手、老实”。这类语素组既不平行,也不周遍,语素组的意义不能从成分和组合关系上得到解释。
- 2. 解释性规则语素组:如“老虎、老鹰、老鼠”等。这类语素组平行但不周遍,其意义可以从成分和组合方式上得到解释,但不能根据这种组合方式无限制地生成新的平行语素组,不可说“老兔、老鸭”。
- 3. 生成性规则语素组:如“老李、老张、老刘”以及“老来、老去、老睡”等。这类语素组既平行且周遍,其意义不仅可以从成分和组合方式上得到解释,还可以根据这种组合方式无限制地生成新的平行语素组。

解释性规则语素组和生成性规则语素组的区分旨在说明,语言的理解过程和生成过程并不一样,生成过程比理解过程要求更严格的条件。理解过程只需要平行的条件,生成过程不仅需要平行条件,还需要周遍条件,因此以上 3 种语素组的语言行为机制并不一样:

	不规则语素组	解释性规则语素组	生成性规则语素组
理解过程	记忆	规则	规则
生成过程	记忆	记忆	规则
语素组	老板	老虎	老刘
判定原则	不平行	平行不周遍	平行周遍
分类	不可解释语符	可解释语符	语符组合

在语言习得过程中,发现生成规则比发现解释规则要困难得多。要听懂言语片断,只需要满足平行条件的解释规则,要说出正确的话语,需要满足平行周遍条件的生成规则。从理解的过程看,解释性规则语素组和生成性规则语素组是一类,不需要记忆;不规则语素组是另一类,需要记忆。从生成过程看,生成性规则语素组是一类,不需要记忆;不规则语素组和解释性规则语素组是另一类,都需要记忆。不规则语素组可以称为不可解释语符,解释性规则语素组可以称为可解释性语符,两者统称为语符。语符库中应该通过标注把不可解释语符和可解释语符区别开。

3.4 平行的性质

以上平行原则的第一个条件要求被替换的部分具有平行特征,这种平行特征可能隐藏在语法、语音、语义等多个层面中,比如前面的“老 X”(老张)中的 X 既有语音特征(单音节),也有语义特征(姓氏)。下面是以语类为平行特征的语素组:

X+的:吃的、写的、买的、用的、大的、小的、公的、金的……

除了少数能够解释的例外,这一格式可以周遍所有的 V、A、D (即通常所说的区别词)语类。但是有些格式仅仅是语类上的平行不一定周遍,以“X 鞋”为例:

N+鞋:皮鞋、铜鞋、冰鞋、水鞋、*气鞋……

V+鞋:跳鞋、跑鞋、拖鞋、*走鞋、*登鞋……

D+鞋(D:区别词):金鞋、银鞋、男鞋、女鞋、童鞋、*公鞋、*母鞋……

以上每一类都有一定的平行实例,但并不周遍。现在我们考虑语义

上的条件,如果我们把 X 的平行特征限制为质料,就可以周遍:

质料语素 + 鞋:布鞋、棉鞋、草鞋、皮鞋、铜鞋、银鞋……

除去 X 表示质料的实例,从其他“X 鞋”实例中还可以抽出一些语义上平行的实例:

运动平行式	项目平行式	组合关系解释
跑鞋	* 步鞋	用于跑步的鞋
跳鞋		用于跳跃的鞋
* 走鞋	* 路鞋	用于走路的鞋
* 踢鞋	球鞋	用于踢球的鞋
* 滑鞋	冰鞋	用于滑冰的鞋
	水鞋	用于下水的鞋

第一列“跑、跳、拖”之间是平行的,这些语素表示的是运动,但不周遍。第二列“球、冰、水”之间是平行的,表示的是和运动项目相关的名物,也不周遍。从两列的关系看,这两组语素和“鞋”组合都是表示跟运动有关系的鞋,第三列给出了平行的语义解释。一个重要的区别是,跑步的鞋要说成“跑鞋”而不说成“步鞋”,但滑冰的鞋却要说成“冰鞋”而不说成“滑鞋”,平行的语义要用不同的表达方式,理由还无法解释。在找到平行周遍条件以前,“跑鞋、跳鞋、拖鞋”都应该看成语符,“冰鞋、水鞋”也应该看成语符,当然,如果“冰鞋”是指用冰做的鞋,比如哈尔滨冰灯节中展出的用冰做的鞋,这时候“冰”满足质料语义特征,“冰鞋”也是规则语素组。“球鞋”的性质后面还要讨论。

平行的第二个条件是组合关系的平行,这是指组合方式和由组合方式表达的组合同义要平行。有些片断的组合关系表面看起来是平行的,但认真对比实例,并不平行。考虑下面语素组:

北京大学、南京大学、武汉大学、上海大学、重庆大学、
昆明大学……

组合方式看起来好像是平行的,前一个成分是城市,后一个成分是“大学”,两者有限定关系,但组合关系的含义实际上并不明确。比如“北京大学”既不是北京所属的大学,也不是在北京的大学。“昆明大学”却是昆明市管辖的一所大学。至少我们现在还不能明确找出“X 大学”组合关系的含义(X 表示地名),所以谈不上平行。因此像“北京大学”这样的片断,都应该存放在语符库中。

并不是因为“北京大学”有特定的指称就必须存放在语符库中,有特定的指称和不满足平行周遍没有必然联系。下面两个片断的指称是相同的:

相对论创始人=爱因斯坦

但“相对论创始人”可以进行平行周遍对比:

相对论创始人、微积分创始人、几何学创始人、物理学
创始人……

因此“相对论创始人”这一类片断并不需要存放在语符库中。与此不同,“爱因斯坦”的指称尽管和“相对论创始人”相同,但必须存放在语符库中。Russell(1905)把“相对论创始人”这类片断称为摹状词,而把“爱因斯坦”这类片断称为专有名词,Russell指出了摹状词和专有名词的区别。从平行周遍的角度看,摹状词和专有名词的一个重要区别就在于摹状词是规则片断,不需要存放在语符库中,而专有名词不是规则片断,需要存放在语符库中。与此相关的一个结论是,我们不应该根据一个片断是否

有特定的指称而断定这个片断是否是语符,而仍然应该根据平行周遍原则。

“北京大学”中“北京”和“大学”的关系不容易定性。在很多情况下组合关系都很难定性,比如下面每一列语素组是什么组合关系,不容易定性:

老张	第一	初一
老李	第二	初二
.....		

但是每一列在组合上的平行关系是容易判定的。这样的实例还可以给出很多。一般地说,判定组合关系是否平行比判定组合关系的性质更容易。这也是我们为什么要从平行周遍入手而不是从定性的角度来区分规则语素组和不规则语素组。

从“老虎、老鹰、老鼠”和“老李、老张、老刘”等实例看,平行性的第三个条件“组合体分布的平行”似乎可以从前两个条件推导出来,也就是说只要满足前两个条件,第三个条件自然就形成了。实际情况并不完全如此,前面的讨论已经初步涉及这个问题,再考虑更多实例:

	被替换项:A类语素	组合关系	组合体的分布
远视	远	偏正	N
弱视	弱	偏正	N
正视	正	偏正	V
傲视	傲	偏正	V
平视	平	偏正	V
轻视	轻	偏正	V,A
重视	重	偏正	V,A
近视	近	偏正	N,A
斜视	斜	偏正	N,V

这里的实例都满足平行的前两个条件,但组合体的分布并不相同。平行性的第3个条件是针对周遍性说的,如果不考虑周遍性,我们可以说“远视、弱视”的分布是平行的,“轻视、重视”的分布是平行的,等等。如果考虑周遍性,我们说“远视、弱视”和“轻视、重视”是不平行的。正是因为我们找不到上述组合体在分布上产生差异的条件,这些实例都需要存放在语符库中。

3.5 周遍的性质

周遍和平行实例的数量没有必然的联系。有些片断的组合可以找到很多平行对比的实例,比如:

盘儿、碗儿、瓶儿、罐儿、刀儿、板儿、门儿……

其特点是指称性语素加“儿”后(或者说儿化后)仍然是指称性语素,有小称的含义。这种平行的实例还可以举出很多,这说明语素“-儿”是组合指数(陈保亚,1999,6)很高的语素,甚至超过了“老X”(老张)的平行实例,但并不周遍,并非所有的指称语素都可以加“儿”,比如“笔、布、筷”等后面就不能加“儿”。哪些指称语素后面可以加“儿”,哪些不可以,周遍条件还没有找到。可见,用平行周遍原则区分语符和规则组合,不在于平行实例的多少,而在于平行实例是否周遍。“老”后边出现的单音节姓氏语素比“儿”前面出现的指称语素要少,但“老X”满足周遍条件,属于生成性规则语素组,而“X儿”不满足周遍条件,属于解释性规则语素组。

从下面两组材料中可以进一步理解周遍的本质:

白纸、白墙、白鞋、白毛、白车……“白菜、”白金、“白铁……

腕儿、腿儿、桌儿、门儿、本儿……‘笔儿、’手儿……

这两组似乎都有不平行的实例,但有根本的区别。第一组不平行的实例“白菜、白金、白铁”是可说的(前面加#号表示可以说),这些实例最初出现时和“白纸、白墙、白鞋、白毛、白车”可能有平行关系,只是后来转义了,除了这些转义的实例,对于任何一个指称语素,只要语义上允许其指称内容有某种颜色,都可以放在语素“白”的后面,和“白纸、白墙、白鞋、白毛、白车”保持平行,比如“白布、白顶、白灯”等。也就是说,“白纸、白墙”这样的语素组格式是平行周遍的,尽管在语言的发展过程中有转义实例出现,但平行周遍集合中减去转义的实例,仍然是平行周遍的。与此不同,第二组的实例“笔儿、脚儿、手儿”和“腕儿、腿儿、桌儿、门儿、本儿”显然是平行的,但前者不可说,不是语义上不允许,也没有其他原因可以解释。“X儿”必须看成是不周遍的。“白纸、白墙、白鞋、白毛、白车”等是生成性规则语素组,“腕儿、腿儿、桌儿、门儿、凳儿”是可解释语符。

平行而不周遍是说在同一平行条件下存在不能说的实例。一般地说,只要我们尚未找出一个平行格式不周遍的特殊理由,或者说无法对不周遍的实例作出充分的解释,这个格式就只能看成平行而不周遍的格式。从语言生成的角度看,这些格式都需要记忆。

平行周遍的情况可以分成有限平行周遍和无限平行周遍。“老X”是有限周遍的情况,因为能出现在X位置上的单音节姓氏语素是有限的。“第X”(X为数词)是无限周遍的情况,因为能出现在X位置上的数字语素或数字语素序列是无限的:

第一、第二、第三、第四、第五……第一百……

“X 的”(X 表示陈述性自由语素或自由语素序列)也是无限周遍的情况:

[写]的、[昨天买]的、[刚看过]的、[准备借给他]的、
[前天上午发现]的……

如果把无限平行周遍语素组储存在语符库中,会违反前面提到的有限原则。如果把有限平行周遍语素组储存在语符库中,可能出现两种情况。当有限平行周遍语素组实例很多时,比如“老张”这类语素组,这时把平行周遍语素组储存在语符库中会违反前面提到的简单性原则。当有限平行周遍语素组实例不多,如“初一、初二……初十”这类语素组,这时平行周遍语素组既可以作为规则语素组处理,也可以储存在语符库中,但放在语符库中除了会增加语符库的项目,同时还要对增加的每个项目进行解释,不完全符合简单原则。无论满足有限平行周遍对比还是无限平行周遍对比,都是规则组合。

用规则来表达的言语片断和文本是无限的,但在实际语言中,只有很少一部分实现了,比如我们听说过“第五”,但可能并没有听说过“第一亿零三十二”。同样,我们听说过“布鞋、皮鞋、草鞋”等言语片断,但可能没有听说过“钢鞋、石鞋、玉鞋、铝鞋”。这并不等于这些片断将来不说,也并不等于这些实例不满足平行周遍。我们需要区分可能出现的片断和已经出现的片断,不能因为有些实例尚未出现就把该实例所在格式当做不规则组合。

既然“钢鞋、石鞋、玉鞋、铝鞋”没有听说过,就表明我们并没有把“X 鞋”周遍完(X 表示质料),我们怎么知道“X 鞋”是规则的?和“X 儿”相比较,“X 鞋”始终没有遇到例外。这里有一个概率抽样的过程,只要没有遇到例外,就可以看成平行周遍的格式。一旦遇到例外,如果例外属于转义的情况,“X 鞋”仍然是有

规则的。如果例外仍然满足平行但不可以说,就要承认不周遍。这种概率抽样的发现过程并不证明平行周遍可以用概率解释,因为平行周遍的本质体现在,对于任何一个从来没有听说过的语素 X,只要可以表示质料,我们就可以生成“X 鞋”。

当然,使用频率很高的规则组合往往容易转义,取得特定的意义,最后成为不规则的组合。和“第十”相比,“第一”出现频率要高得多,现在“第一”已经取得了“首先”的意义,可以和“其次”、“最后”等相互照应。从工程技术的角度看,高频率的有规则语素组合有时候也可以放在语符库中,这样有利于在文本中自动切分语符,但缺点是使得语符库收录标准不统一,规则语素组和不规则语素组混在一起,记忆、理解、生成三种机制的差异得不到反映。当自动切分需要高度精确时,频率标准会成为障碍。另外,规则语素组频率达到多少就可以储存在语符库中,没有统一的标准,比如“布什总统”是一个出现频率很高的片断,是否可以储存在语符库中? 会引出很多争议。为了避免这些缺点,还是应该严格坚持平行周遍原则,只有当规则语素组发生转义时,才需要把转义的项目储存在语符库中,比如具有“首先”义项的“第一”就可以存放在语符库中。

平行周遍有一个发现过程。有时候两个语素的组合是否有规则,跟我们是否已经找到规则有关系。当两个语素的组合规则还没有找到的时候,这两个语素是语符,当两个语素的组合规则找到的时候,就是规则组合。比如前面讨论的“X 鞋”这种组合中,当 X 表示质料这个平行条件被找到后,“X 鞋”就是平行周遍的,是规则组合,但是在我们发现这个平行周遍条件以前,“皮鞋、布鞋、草鞋”等语素组都应该储存在语符库中。语符库不是先于组合规则的分析而形成的。语素的提取不需要组合规则的信息,语符的提取却需要组合规则的信息。语符的提取和组合规则的提取是同一项工作的两个方面,随着组合规则研究的深入展开,提取到的规则越来越多,语符的数量越来越少,这是

一个动态过程。有时候从直觉上能够感到规则的存在,但平行周遍条件并没有给出来,这时的实例都要作为语符处理,因为这时我们没有办法告诉第二语言学习者规则在哪里,因此必须让他记住实例。

平行周遍发现过程的另一个方面表现在规则的统一解释上。“布鞋”和“白鞋”这两类情况都满足平行周遍对比,“布鞋”类的平行周遍条件是“质料”,“白鞋”类的平行周遍条件是性质。这两类可以进一步统一起来,因为“质料”和“性质”可以统一看成属性,统一的平行周遍规则是:

X 表示属性,鞋具有 X 的属性。

即使统一的解释规则没有找到,“布鞋”这一类组合和“白鞋”这一类组合仍然分别作为规则组合处理。有无规则和是什么规则属于两种不同的情况,前者是对无规则而言的,后者是对规则描写的简单程度而言的。

3.6 多个语素序列的判定

前面我们说“冰鞋、球鞋”这样的平行实例是不周遍的,这是从前项的替换观察平行周遍情况。如果从后项看,“球鞋”可以进行平行周遍替换:

球鞋、球衣、球裤、球帽、球袜、……

平行规则是:

“球 X”中,X 表示运动的服装。

因此,“球鞋、球衣”等可以看成规则语素组。但“冰鞋”从后项看,似乎也找不到平行周遍的实例,所以“冰鞋”是不规则语素组。一般地说,有些组合满足双项平行周遍对比,即前后两项都可以进行平行周遍对比,有些组合只有一项可以进行平行周遍对比。比如:

双项平行 周遍对比		单项平行 周遍对比 1		单项平行 周遍对比 2	
布鞋	布鞋	球鞋	球鞋	同学们	同学们
草鞋	布袜	冰鞋	球衣	士兵们	?
皮鞋	布帽	雪鞋	球帽	老师们	
.....					

第 1 列的“布鞋”前后两项都可以作平行周遍替换。第 2 列的“球鞋”只是后项可以作平行周遍替换,前项不周遍。第 3 列“同学们”只是前项可以进行周遍替换,后项没有平行实例可比。

两项中只要有一项可以进行平行周遍对比,该语素组就是规则语素组,这为判定规则语素组提供了程序化的可能。一般地说,对任何一个由 n 个语素构成的片断,只要其中任何一个语素或语素组合可以作平行周遍对比,这个由 n 个语素构成的片断就不是语符而是规则组合。我们把这个结论称为平行周遍的多序列判定原则。这个原则的价值在于,当语素组某个位置上成分的平行关系不好确定时,可以考察其他位置上成分的平行周遍情况。比如从语感上看,“两个人”是规则组合,但我们目前还不知道满足什么条件的名物语素或名物语素组合可以出现在“两个”后面,只能举例说明:

两个 X: 两个人、两个梨、两个头、` 两个猪、` 两个书.....

由于 X 出现在“两个”后面的条件没有找到,不可能通过 X 实例的平行情况来断定“两个 X”是否是规则组合,但是这个组合中

所包含的“两”是可以作以下平行周遍对比的：

两个人、三个人、四个人、五个人、六个人、七个人……

所以“两个 X”作为规则组合的地位仍然可以通过平行周遍对比确定。按照多序列判定原则，可以判定“将他的军、理一次发”都是规则组合，“将军、理发”都不是规则组合。但多序列原则不能确定语素组合的层次。

在多序列判定原则基础上，不连续语素组或其他更复杂的语素组是否是规则组合也容易得到判定。比如“越走越快”：

“走”被替换：越走越快、越做越快、越吃越快、越看越快……

“快”被替换：越走越快、越走越慢、越走越累、越走越远……

在“越 V 越 A”中，V 和 A 都可以被替换，替换后满足平行周遍，剩下的“越……越……”不能再进行平行周遍对比，也应该看成语符。由于“越……越……”是不连续的，可以称为不连续语符。通常所说的固定格式或固定搭配，都不能进行平行周遍对比，如果这些固定格式或固定搭配是不连续的，性质上都应当属于不连续语符。

与多序列判定原则相关的一个问题是，如果某个片断可以进行平行周遍对比，包含这个片断的更大片断是否也可以进行平行周遍对比？回答是否定的。“狗肉”可以进行平行周遍对比，但包含“狗肉”的片断“挂羊头卖狗肉”就不能进行平行周遍对比。类似“挂羊头卖狗肉”这样的片断应该存放在语符库中。语符库是存放不规则语素组的，凡是不满足平行周遍的片断，都应该存放在语符库中，不论组合的长短。

3.7 语素库、语符库和词库

Bloomfield(1926,1933)的单说论和陆志韦(1957)的扩展法是建立在区分自由和黏着这一对概念上的,从前面的分析可以看出,平行周遍和自由、黏着没有严格的对应关系,下面把它们的相互关系列表如下:

语素组实例	前项	后项	平行周遍情况
布鞋、草鞋、纸鞋	自由	自由	平行周遍
跑鞋、跳鞋、“滑鞋	自由	自由	平行不周遍
白菜、花菜、青菜	自由	自由	不平行不周遍
老李、老张、老刘	黏着	黏着	平行周遍
老虎、老鼠、“老兔	黏着	黏着	平行不周遍
老师、老财、老板	黏着	黏着	不平行不周遍
第一、第二、第三	黏着	自由	平行周遍
泛读、泛听、“泛写	黏着	自由	平行不周遍
老家、老气、老手	黏着	自由	不平行不周遍
走了、跑了、来了	自由	黏着	平行周遍
本儿、书儿、“笔儿	自由	黏着	平行不周遍
学籍、学业、学费	自由	黏着	不平行不周遍

以上凡是平行但不周遍的,只列出一个规则上平行但不可以说的例子。容易看出,判定规则语素组和不规则语素组不能根据自由和黏着的标准,即两个自由的语素组合后可以是不规则语素组,两个不自由的语素组合后可以是规则语素组,一个语素黏着而另一个语素自由,可以是规则语素组,也可以是不规则语素组。当然自由的概念在平行特征中有时会用到,比如“老来、老去、老睡”中的“来、去、睡”都是自由的 V 类语素或 V 类语素组,这时候自由的概念仅仅是作为一个平行特征来使用,并不证明自由和规则之间必然有对应。自由和黏着的价值体现在语法

分析的其他方面。

很多语言学家区分句法和词法两个部分,词法研究词的内部构造,以语素为基本单位,句法研究句子的内部构造,以词为基本单位。这个区分是以自由和黏着为标准的,自由语素组或可以扩展的语素组是句法研究的对象,否则是词法研究的对象。下面把规则语素组和不规则语素组的区分跟句法和词法的区分作一个比较:

句法	词 法			
句法	形态(构形)	附加构词 1	附加构词 2	复合构词
写书	写了	老李	桌子	桌布
白车	写过	第三	石头	石器
染发	染的	初三	老虎	理发
老写	写着	阿哥	花儿	花瓶
规则语素组			不规则语素组	

上面关于词法的再分类只是为了比较的方便,在术语上不一定准确。容易看出,句法涉及的是规则语素组,词法中包括了规则语素组,也包括了不规则语素组。对句法和词法的区分可以提出下面一些问题:

1. 附加构词 1(“老李”)和附加构词 2(“桌子”)从平行周遍看是完全不同的两类,一类是有规则的,一类是没有规则的,都被当做词法中的附加构词处理。

2. 附加构词 1 中的“第三”,组合模式是“第 X”,平行实例是无限的,都作为词来处理,违反了单位的有限原则。

3. 句法和词法的区分并没有解释可推导性,也不是为了区分规则语素组和不规则语素组。

4. 目前没有严格的标准能够把句法和词法区分开,前面讨论单说、扩展、转换等标准时,已经说明这一点。

如果我们严格按照规则来处理组合层面,可以分出语法和词汇。语法包括句法和构形。语法的基本单位是语符和词。词汇包括构词、词义解释以及词的特殊组合性质(比如题元)。

语符可以作以下分类,并可以和目前的词本位、字本位(徐通锵,1997)作一个比较:

语符本位	实例	词本位	字本位
不自 由语符	实语素:鸭、翅	构词成分	字
	词头:第、老		
	后缀:化、性		
	量词:个、只	词	
	虚词:了、过、着、的、很、从		
	区别词:金、银、男、女		
自由语符	实词:人、手、走、吃、大、多		

一般地说,一个词通常是一个语符(如“人”)或语符组(如“鸭翅”),但一个语符不一定是一个词,只有自由语符通常是一个实词。在后面的讨论中,我们将区分规则单位库,一种是语符库,按照语符标准存放单位,一种是词库,即通常所说的 lexicon。由于不同的学者关于词库中存放单位的标准不一致,不太容易展开讨论,所以我们后面会更多地用语符库的概念。尽管不同的学者对词有不同的理解,但从以上比较容易看出,词库是语符库的一个子库,只是这个子库的确定标准不够严格,或没有达成共识。当然“自由”的概念仍然是相当重要的,因为只有自由语符或实词才可以单说、移动和充当句子的直接成分和论元。有时候为了术语的统一,我们也用词库这个术语,但如果没有特别说明,都是指语符库。

语素不能直接作句法单位,但语素的地位仍然很重要。判定规则语素组和不规则语素组要以语素为前提。后面我们要讨论,音系研究也要以语素音形为前提。因此,我们需要在符号层面建立两种性质不同的单位库:语素库和语符库。

徐通锵(1997)关于字的定义和语素不完全相同,不同的学者有不同的理解,后面暂时不展开字本位的讨论,当涉及具体问题的时候会加以分析。

3.8 派生、转换和语符库的关系

Chomsky(1965,p164)开始考虑词库和句法的关系,其中一个重要思想是,如果词库中包含的次范畴特征较少,选择特征多,就更有价值。

Furthermore, let us make the empirical assumption that a grammar is more highly valued if the lexical entries contain few positively specified strict subcategorization features and many positively specified selectional features.

Chomsky(1965,4. 2. 3,p184)根据名物化转换(nominalization transformations)把 *destruction*, *refusal* 这样的词处理成转换的结果,因此这些词不需要进入词库:

Clearly, the words *destruction*, *refusal*, etc., will not be entered in the lexicon as such. Rather, *destroy* and *refuse* will be entered in the lexicon with a feature specification that determines the phonetic form they will assume (by later phonological rules) when they appear in nominalized sentences.

参考实例分析容易看出,Chomsky 把名物化这个过程形式化为:

$$N \rightarrow \text{nom} VP$$

Chomsky(1965, 4. 2. 3, p185)进一步认为:

In any event, phonological rules will determine that *nom~destroy* becomes *destruction* and that *nom~refuse* becomes *refusal*, and so on.

即根据具体的语音规则,就可以得到不同的动词的名物化转换形式:

nom~destroy→*destruction*
nom~refuse→*refusal*

Chomsky(1965)认为其代表例句 *sincerity may frighten the boy* 中的 *sincerity* 也不应该进入词库,而应该处理成通过对 *sincere* 施加转换规则形成的。这里只代表了 Chomsky(1965)标准理论的思想,他后来的思想有所不同。但这种思想现在仍然是很多人的主张,有必要加以分析。我们认为,名物化的平行周遍规则还没有找到,以-ion 为例:

attraction
seduction
imitation

在没有找到平行周遍条件之前,要处理成转换生成规则是有问题的。如果 *destruction* 不收入词库, *destroy* 要标注下面的属性:加-ion 变成名词。在没有找到周遍条件以前,每个可以带-ion的都要标注,有 n 个这类形式,就有 n 个标注,跟直接把带-ion的形式存放在词库里没有区别。

当然,不是说不可以把部分派生词从词库中提出来,关键是

要区分平行不周遍条件和平行周遍条件。可以说生成语法没有解决这个问题。从 Chomsky(1965, p186) 分析的实例可以进一步看出当时对规则的性质认识不清楚。例如:

horror	horrid	horrify
terror	* terrid	terrify
candor	candid	* candify

Chomsky 认为这些项目(item)都应该进入词库,因为这些都属于准能产过程(quasi-productive process)。但我们认为前面 Chomsky 关于 destroy 的派生过程和这里的情况是类似的,也是不能产的,从根本上说,也是不平行周遍的:

destroy	* destruct	destructive	destruction
* constroy	construct	constructive	construction
* obstroy	obstruct	obstructive	obstruction
* instroy	instruct	instructive	instruction

Chomsky(1965, p187)还提到或许可以把这些未发生的形式看做规则模式中的一种空格,后面我们会讨论把这些例子作为规则模式所遇到的困难。

关于能产性和转换问题曾经有过比较深入的讨论,分成词汇派(lexicalist position)和转换派(transformational position)。从 Chomsky 上面的分析可以看出,当时 Chomsky 已经提出了词汇派观点(lexicalist position),并且有比较明确的态度(Chomsky, 1965, p219-220)。词汇派的基本观点是主张把上述不能产的或准能产的派生形式都存放在词库中。转换派的基本观点是把这些准能产的形式看成是转换的结果。转换派的代表有 Lees(1960), Lakoff(1965)。也有些学者采取中立的态度。

Chomsky(1970)区分了 refusal 和 refusing 两种现象,认为 refusal 是派生的结果,而 refusing 才是转换的结果,Chomsky 区分这两类情况有三个方面的理由:

1. V+ing 中 V 不受限制,但是 V+al 中 V 要受限制(能产性问题)
2. 从 V 到 V+ing 语义关系是规则的,但从 V 到 V+al 的语义关系并不规则
3. V-al 具有一般名词的活动方式,V-ing 不是这样。

比较起来,Chomsky 的词汇派观点更能说明词库和规则的关系。V-al 受到限制,这是能产性的问题。但是为什么会有转换派的支持者,我们认为是因为能产性问题的复杂性没有完全弄清楚。Chomsky 在区别词性变化和派生的基础上来讨论可转换和不可转换问题,说明当时还没有完全弄清规则和不规则的本质。比如,Chomsky(1970,p187—189)在讨论能产性问题时,列举了以下几种情况,其中动名词的活动方式是能产的,但派生名词的活动方式受到限制:

John is easy (difficult) to please

John's being eager to please

John's being easy (difficult) to please

John's being eager to please

* John's easiness (difficulty) to please

John's eagerness to please

Chomsky(1970,p191)认为动名词的转换成立而派生名词的转

换不成立,是因为深层结构不一样。

John is eager (for us) to please
(for us) to please John is easy

这种解释是没有问题的。再比较更多的例子:

- 6 a. John is easy (difficult) to please.
- b. John is certain (likely) to win the prize.
- c. John amused (interested) the children with his stories.

- 7 a. John's being easy (difficult) to please
- b. John's being certain (likely) to win the prize
- c. John's amusing (interesting) the children with his stories

- 8 a. * John's easiness (difficulty) to please
- b. * John's certainty (likelihood) to win the prize
- c. * John's amusement (interest) of the children with his stories

- 9 a. John's eagerness to please
- b. John's certainty that Bill will win the prize
- c. John's amusement at (interest in) the children's antics

- 10 a. John's being eager to please
- b. John's being certain that Bill will win the prize
- c. John's being amused at (interested in) the children's antics

我们认为这里的实例正好证明了一个谓词带上补语(complement)后有哪些不同的转换方式是由谓词的个性决定的,所以词库必须记录该谓词的个性。但由此来证明动名词是转换的结果而派生是词库信息,却存在方法论上的问题。我们认为,正是因为我们能够从深层结构上指出为什么上述例子中含 easiness 的不能说而含 eager 的能说,因此可以解释转换条件,也就是说,eager 由于在深层结构上可以带一个句子作补语,所以上述转换是可以的。以上动名词和派生名词之间的真正区别是,派生名词的派生规则并不是周遍的。后面我将从平行周遍的角度来认识转换的规则和不规则问题。一般地说,一种转换的周遍性,和动名词并没有一致关系,一种转换的不周遍性,和派生名词也没有一致关系。汉语动词的主动形式和把字句形式的关系被看成是句法现象,但从主动形式到把字句形式的周遍条件现在还没有找到。汉语的数词和“第 X”的关系被看成是派生现象,但周遍条件却是明确的。英语中数词的派生形式“X-th”也有周遍条件。

3.9 短语的语类推导

前面说到 Harris(1946)从语素推导话语的模型由于把句法单位建立在语素上会遇到困难。如果把句法单位建立在有生成能力的单位上,应该是可行的。詹卫东(2000)曾经在词的基础上展开过汉语语类推导,由于词是有生成能力的单位,是可行的。现在我们把句法单位建立在语符上,一些从词的角度看被作为短语词的组合,也可以实现从单位到话语的预测。更准确地说,可以实现从单位的语类到组合的语类的预测。如果语符库中充分收集了单位和不规则组合,那么其他组合就可以作为规则组合来处理,这时从单位分布预测规则短语或规则组合的分布就可以实现。所以找出单位是生成规则短语的必要基础,

否则就会生成不合法的组合。提取单位在语言研究中有重要意义。

一般地说,如果我们充分区分了规则组合和不规则组合,对语符的类作了充分描写,就可以进一步描写规则组合的类。现在我们以汉语为例来考虑在单位分布的基础上预测规则组合的分布。先考虑 N 前面连接 A、V、N、D、Prep、Xd 的情况。

$A+N=NP$:

小房子、大房子、宽房子、红房子、蓝房子、白房子、绿房子、高房子、矮房子

$N+N=NP$:

石头房子、铁房子、木头房子、水泥房子、砖房子、玻璃房子

$D+N=NP$:

大型房子、小型房子、老式房子、新式房子

$Xd+N=NP$:

漂亮的房子、干净的房子、宽敞的房子

出售的房子、装修的房子、出租的房子

木头的房子、水泥的房子、玻璃的房子

$V+N=VP$:

修房子、盖房子、买房子、拆房子

$Prep+N=PP$:

在学校、在书房、在车间、在澡堂

要进行语类推导,必须要区分规则组合和不规则组合,这是

一个必要条件。当然不是充分条件。因为在单位的语类已经确定的条件下,组合有时会有不同的类,这方面的条件还没有完全弄清楚。比如:V+N

NP/VP: 烧土豆、炒土豆、烤土豆、煮土豆、蒸土豆

学习材料、出版书籍、打印稿件

VP: 买土豆、买材料、送书籍、送稿件

烧房子、烤衣服

以上 V+N 形成的歧义,不仅和 V 的语类有关系,和 N 的小类也有关系,甚至和语义类也有关系。比如烹调类的动词就可以造成 V+N 的歧义。部分 VN 的语类推导条件可以通过细分词类来确定。也可以有另一种方式来处理部分情况。一般地说,只要双音节的 V 能起到分类的作用,VN 都有作偏正结构或者 NP 的可能。也就是说,双音节 V 形成的 VN 本来就有两种分布,是哪一种取决于语义条件。比如下面两组动词都相同,由于“材料”和相应的动词组合可以允许分类,所以有 NP 的可能,“英语”和相应的动词组合不存在分类的可能,所以没有 NP 的可能,这也说明 V 能否作修饰语不完全取决于动词分布。

VP/NP: 学习材料、调查材料、分析材料、研究材料、准备材料

VP: 学习英语、调查英语、分析英语、研究英语、准备英语

不过这只是一般的认知结构,如果考虑更多的例子,V 和“英语”类名词组合也可以是 NP:

补习英语、参考英语、详解英语、翻译英语

即只要语境允许这样分类,这里的组合仍然有 NP 的可能。

至于单音节动词,有些条件可以找到,比如“炒、煮、烤、炖、烧”等是和烹调有关系的,VN 可以是 NP。至于其他没有找到条件的,可以继续在意符库中标注名动词来完成语类推导。关于这方面的复杂情况,邱立坤(2004)有进一步的调查分析。

以上语符的语类推导出现了 NP 和 VP 等组合,我们再来看这些组合前面继续叠加的情况:

$A + NP = NP$:

小[红房子]、小[蓝房子]、小[白房子]、小[矮房子]

新鲜[炒菜]

$N + NP = NP$:

木头[小房子]、玻璃[小房子]

四川[烧菜]、东北[炖菜]

$D + NP = NP$:

旧式[小房子]、旧式[矮房子]

旧式[木头房子]、旧式[水泥房子]

$X_d + [A + N]$:漂亮的[新房子]、调查的[新材料]

$X_d + [N + N]$:漂亮的[木头房子]、出租的[木头房子]

$X_d + [V + N]$:丰富的[英语材料]、出售的[英语材料]

$X_d + [X_d + N]$:宽敞的[漂亮的房子]、出售的[崭新的房子]

$X_d + [D + N]$:宽敞的[旧式房子]、出租的[旧式的房子]

$V + NP = VP$:

$V + [A + N]$:学习[新材料]、调查[新材料]、打印[新材料]

V+[N+N]: 打印[英语材料]、出售[木头房子]

V+[V+N]: 打印[学习材料]、出售[打印材料]

V+[Xd+N]: 出售[危险的房子]、打印[学习的材料]

V+[D+N]: 出租[旧式房子]、购买[旧式房子]

Prep+NP=PP:

Prep+[A+N]: 在小房间、在大澡堂、

Prep+[N+N]: 在木头房间、在学校澡堂

Prep+[V+N]: 在补习学校、在打印车间

Prep+[Xd+N]: 在弟弟的房间、在住的地方

Prep+[D+N]: 在新式澡堂、在大型车间

NP 被 V 叠加后构成 VP, 被 Prep 叠加后是 PP, 被其他语类的语符叠加后仍然构成 NP, 整个叠加后的组合在语类上并没有越出 VP、PP 和 NP。前面已经讨论过, N 被单个语符从前面叠加后也都是 VP、PP 或 NP, 我们说 NP 从前面叠加已经进入递归状态, 三个语符构成的 NP 前面再继续叠加语符也仍然是 NP、PP 或 VP 三种情况, 不会出现新的语类。我们对其他语符组合的调查也显示, 语类叠加在有限步骤内会进入递归状态。如果出现例外, 基本上都是属于习语类型。以上语类连接、叠加的语类推导规律可以概括为:

在规则组合中:

1. 语类的连接、叠加所得到的组合的语类是可以推导的。
2. 语类的连接、叠加在有限步骤中可以到达递归状态, 因此语类是有限的。

以上我们只是在连接、叠加中讨论语类推导, 由于扩展替换所得到的组合是叠加所得到的组合的子集或者说特殊情况, 所以以

上语类推导规律对替换也是有效的。

Harris(1946)没有讨论直接成分或层次问题,其实,Harris的语素话语模型可以直接延伸到层次,因为Harris要求,构成话语的语素或语素序列都必须是可以归类的话语才是符合语法的。由上面给出的推导过程可以看出,只要语素序列从小到大都可以归类,这种层次就已经满足了语类条件(当然最后的层次确定还要满足语义条件)。后来Wells的直接成分理论正是建立在这个基础上的。具体表现在两个方面:

1. 相同的语素序列有些可以进行语类规约,有些不可以,可以规约的就形成一种层次。比如:

$ad + aj + 的 + n$

$((\text{很聪明})\text{的})\text{孩子} \rightarrow ((ad + aj) + 的) + n$

$\rightarrow (aj + 的) + n \rightarrow n + n$

$\text{很}((\text{聪明的})\text{孩子}) \rightarrow ad + ((aj + 的) + n)$

$\rightarrow ad + (n + n) \rightarrow ad + n \rightarrow ?$

2. 相同的语素序列可以作不同的规约,因此有不同的层次。比如:

$[_{NPd} \text{三个研究所的}] \text{老师}$

三个 $[_{NP} \text{研究所的老师}]$

可见层次必须符合替换或语类的限制。对于第一种情况来说,只有可替换的才是正常的语素序列。对于第二种情况来说,有两种替换就有两种层次。这说明满足分布类的替换要求是层次的必要条件。

以上语类推导是从单位的语类推导组合的语类,但单位的

语类不完全能推导出结构关系,后面我们证明结构关系是初始概念时,要继续讨论这个问题。另外,有些推导中相同的语类成分可能推导出不同语类结构,这类现象往往都可以通过细分语类得到解释,或者有特殊的条件,这里不展开。

思考问题

1. 如何给平行周遍原则一个更好的定义?
2. 汉语的句法单位能否建立在字的基础上?

4. 替换、叠加和成分线性特征

从结构语言学到转换生成语法的一个根本转变是分出句法的基础部分和转换部分,这也是转换生成语法后来一直保持的一个基本思路。区分基础部分和转换部分的核心问题之一是如何确定深层结构的性质、范围或界限,这就需要深入认识线性结构的性质。下面我们仍然从替换、叠加和线性特征入手展开问题的讨论,我们所说的单位都是指上一章谈到的语符。

4.1 自动机理论与自然语言的成分线性特征

考虑三个单位的线性关系:

ABC

我们可以分出两种线性关系。一种线性关系是单位线性关系:

$A+B+C$

还有一种关系是成分线性关系:

$(A+B)+C$ 或 $A+(B+C)$

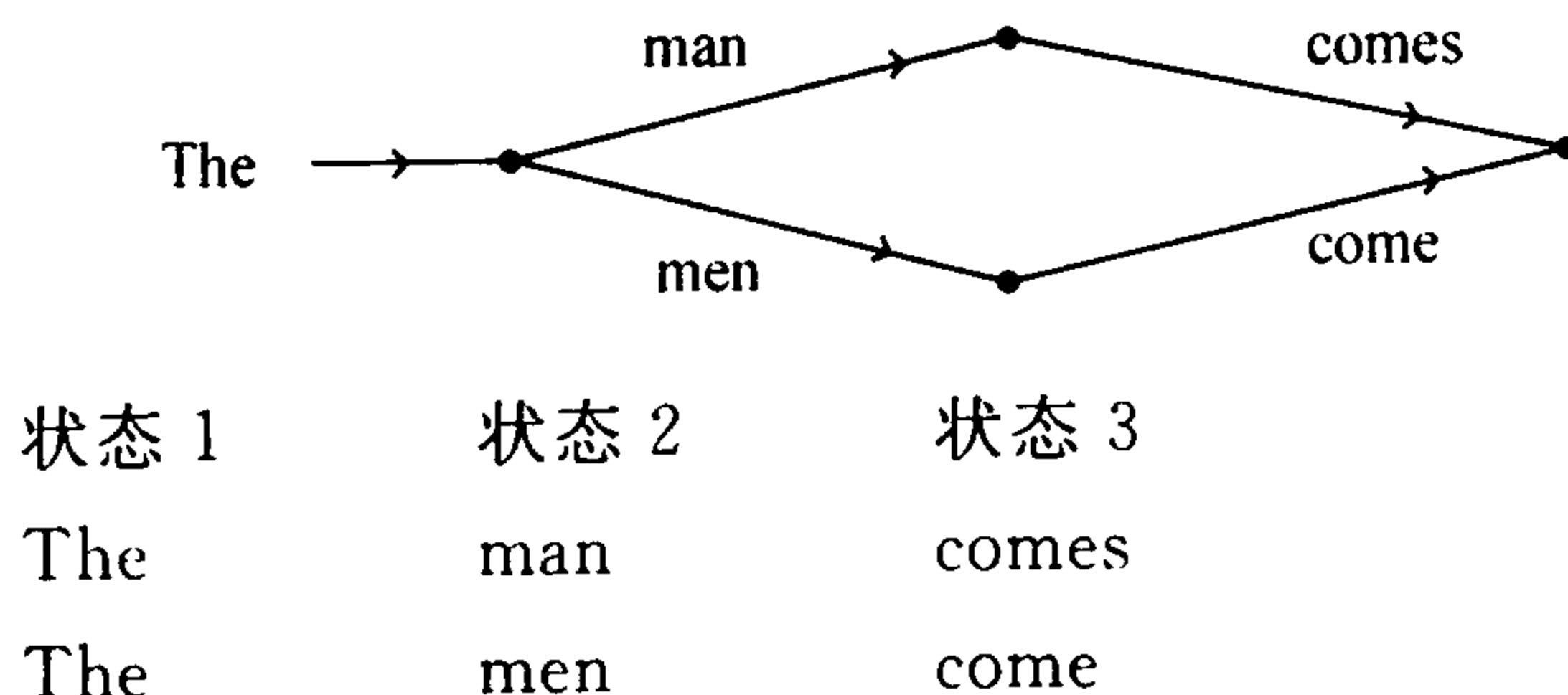
第一种线性关系可以说是真正意义上的单位线性组合关系,就像项链上各个环节的组合关系。第二种线性关系则说是直接成分的线性关系。当 Saussure(1916)提出符号线性关系的时候,

并没有明确说明符号的线性关系是这两种关系中的哪一种。从 Saussure 的行文表述中可以发现,他的线性关系相当于单位线性组合关系。Bloomfield(1933)提出了直接成分的理论。我们认为直接成分的存在说明还有另一种线性关系,即直接成分线性关系,不过 Saussure 的单位线性是否存在,仍然是一个重要的认识论问题。我们认为,从整个符号学的角度看,这两种线性关系都是最基本的线性关系,现在要讨论的问题是,自然语言是否也存在这两种线性关系。

下面我们将论证,在自然语言中,只要是单位的线性组合(自然语言还包括非线性组合,不属于要讨论的范围),符号的线性关系都是直接成分的线性关系,左线性、右线性也是直接成分线性关系的特殊情况。自然语言中不存在符号的单位线性关系。两个单位的组合是直接成分线性关系的特殊情况。

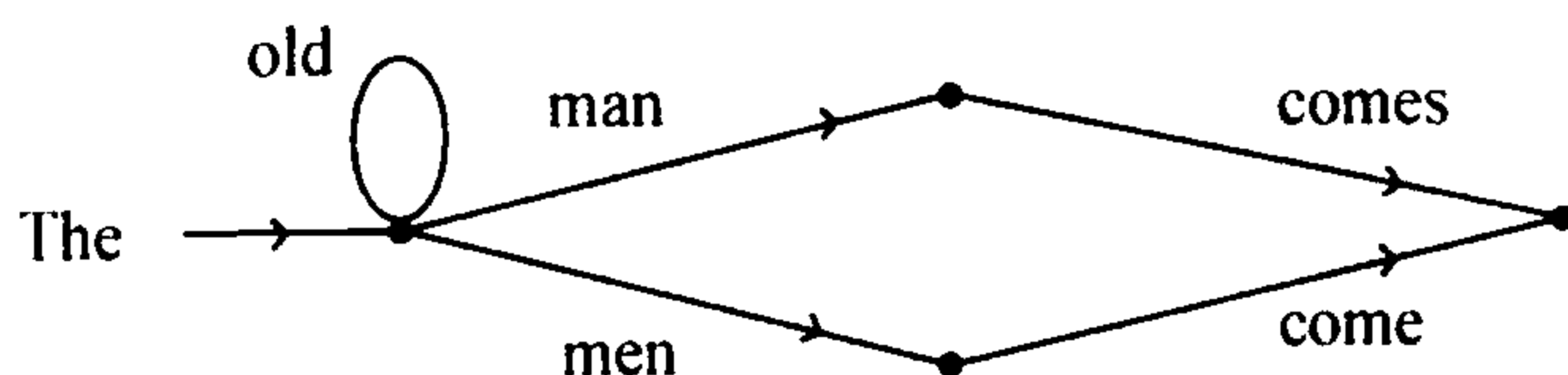
Chomsky(1957)讨论了有限状态语法、短语结构语法和转换语法三种模型,对很多重要理论问题作了深刻的分析。现在我们先在 Chomsky 的基础上讨论有限状态语法的线性性质,因为这个模型是认识自然语言符号组合的线条性的关键。

Chomsky(1957, 3.1)认为语法的一个要求就是要满足有限性(finite),因此语法不应该仅仅是语素序列或词序列的清单(a list of all morpheme (or word) sequences),因为语素序列是无限的。他认为通讯数学模型中的有限状态自动机(finite state machine)满足从有限规则生成无限话语序列这个要求。这种装置从初始状态,经过有限的中间状态,达到终止状态。Chomsky 把有限状态自动机生成的序列叫有限状态语言(finite state language),并从语法的角度把有限状态自动机的规则称为有限状态语法(finite state grammar)。借用 C. E. Shannon 和 W. Weaver 所给出的通讯数学模型(The Mathematical Theory of Communication, 1949),Chomsky 描写了英语的句子,分析了有限状态语法的性质。观察下面实例:



如果把状态 1 看成初始状态,下一个状态(即状态 2)的选词是由上一个状态(即状态 1)决定的,即从状态 1 的 the 出发可以有状态 2 的 man 或 men 两个选择。状态 3 的选词又是由状态 2 决定的,即如果状态 2 选择了 man,状态 3 就只能选择 comes,如果状态 2 选择了 men,状态 3 就只能选择 come。因此有限状态语法是靠“从左到右”在每个状态选词来生成句子(Chomsky, 1957, 3, p14)。说话人从第一个状态开始,生成句子的第一个词,然后转到第二个状态,这个状态限制了第二个词的选择。如此下去,每个状态都代表了一个下一步选词的限制。

上面描述 the man comes 的有限状态语法生成出的句子是有限的,不满足生成性的第二个思想。因此还必须采取加封闭圈或回路(loop)的办法,使有限状态语法有无限的生成能力。所谓加封闭圈,就是允许从某个状态返回到前一个状态或允许某个状态返回到状态自身。以状态返回到自身为例:



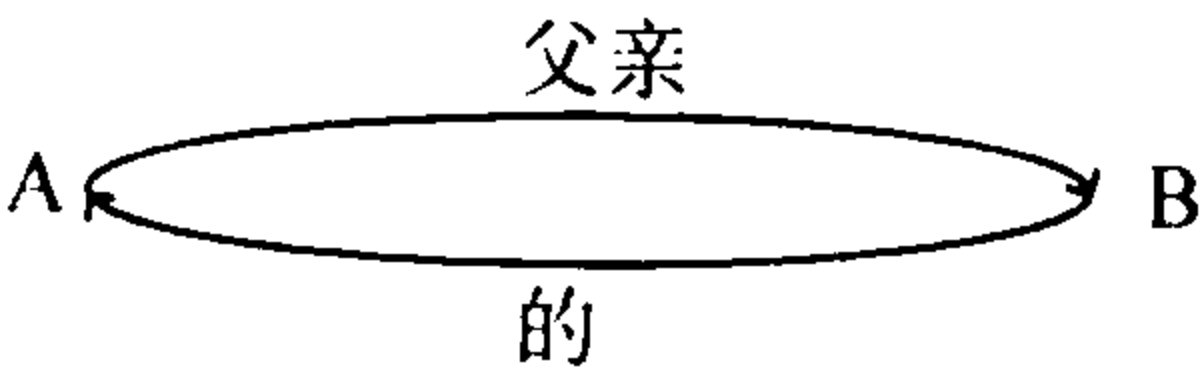
这一封闭圈每循环一次,就生成一个 old,循环 n 次就可以生成 n 个 old,这就保证了这一有限状态自动机所生成的句子是无限的:

The old man comes	The old men come
The old old man comes	The old old men come
The old old old man comes	The old old old men come

Chomsky 讨论的循环实例是同一个状态或节点循环。循环也可以在不同的状态间进行。让我们再通过汉语的一个实例来认识不同状态间相互返回的情况：

父亲的
父亲的父亲
父亲的父亲的
父亲的父亲的父亲
.....

由“父亲”和“的”两个词不断叠加可以构成的组合是无限的，我们设计一个 A、B 两个状态的有限状态自动机或有限状态语法，就可以生成这些无限的句子：



这台自动机从 A 状态开始，A 状态选择“父亲”进入 B 状态，B 状态选择“的”又可以进入 A 状态，A 状态选择“父亲”又可以进入 B 状态。根据需要，可以生成任何长度的“父亲”和“的”的组合。

现在来讨论最关键的线性问题。Chomsky(1957)论证说有限状态语法可以描写自然语言的一部分句子，并认为有限状态语法是建立在 Markov 过程(markov process)上的。Markov 过程是一类重要的随机过程，其原始模型为 Markov 链，由俄国数学家 Markov 于 20 世纪初提出。Markov 和后来的数学家没有专门从语言学或符号学的角度考虑 Markov 链的线性性质，

Chomsky 也没有。但从 Markov 链所描述的现象看,实际上 Markov 链既可以是单位线性的,也可以是成分线性的,所以 Chomsky 并没有说明有限状态语法是建立在单位线性 Markov 过程上的,还是成分线性 Markov 过程上的。我们下面要论证,有限状态语法并不是单位线性 Markov 过程,而是成分线性 Markov 过程。单位线性 Markov 过程不可能描写自然语言句子。自然语言的线性组合不是相邻单位的线性组合,而是相邻直接成分的线性组合。

Markov 过程的基本思想是:一个事件在已知的当前状态下,它的下一个状态不依赖于它以往的状态,只与当前状态有关。用数学模型表示就是:

$$q(t) = f(q(t-1), s(t)) \quad (t \text{ 大于或等于 } 1)$$

其中 $q(t)$ 表示事件在时间 t 的状态, $q(t-1)$ 表示时间 $t-1$ 时的状态, $s(t)$ 表示时间 t 的输入信号。怎样理解 Markov 的当前状态是认识自然语言线性性质的关键。如果我们直接从语言的单位线条性中来认识 Markov 过程的本质,就会得到下面一种描述方式:

状态 1	状态 2	状态 3
the	man	comes

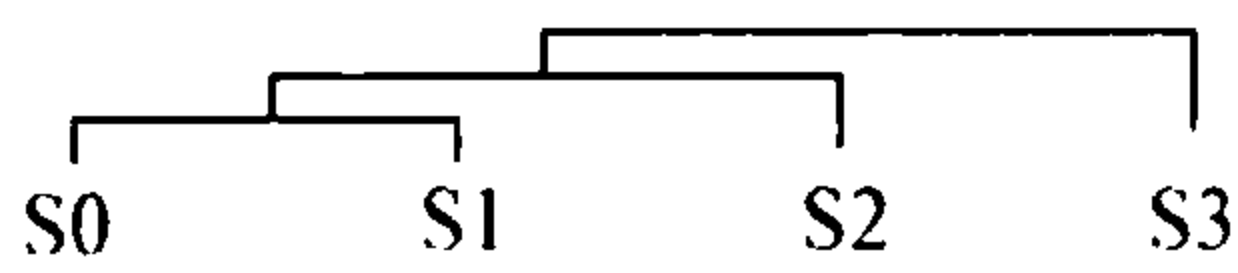
如果从成分线性来理解 Markov 过程,就有:

状态 1	状态 2	状态 3
the	the man	the man comes

一般语符序列的成分线性 Markov 过程就成了:

$$[[[S_0, S_1,]S_2,]S_3,]\cdots$$

其中 S_i (i 取自然数)每一对方括弧都代表一个状态,其更直观的树形图是:



我们在前面讨论过自然语言线性生成的两种方式,一种是叠加,一种是替换,由这两种方式生成的组合都是直接成分的组合,都是有层次的。也就是说,在符号构成的线条中,除了初始组合以外,只有直接成分的组合,并不存在单位的直接组合。因此我们说,除了语符首次连接构成的基式(即两个单位的组合)可以是单位线性组合,在基式上进行的再次叠加或替换而形成的组合,都不能用单位线性组合解释。让我们再回到 Chomsky(1957)提到的由 3 个词构成的句子 the man comes,当时 Chomsky 认为这个句子可以用有限状态语法解释,因此也可以用 Markov 链解释。我们认为这个过程尽管可以用有限状态自动机生成,但不可以用单位线性原则加以描写,因为这个实例的线性关系是这样的:

$$[the\ man]\ [comes]$$

即 comes 并不是和前接的 man 发生组合关系,而是和 [the man]发生组合关系。

这里需要深究“当前状态”的含义。Chomsky(1957)的有限状态语法中的当前状态必须以对 the man 的规约为前提,即把 the man 看成一个直接成分,所以 Chomsky 有限状态语法、Markov 过程、有限状态自动机等 在识别 the man comes 时,都已经涉及直接成分。后面我们将通过替换的性质进一步说明,

在自然语言中,语言符号的线性组合并不存在单位线性组合,只存在直接成分线性组合。

有限状态自动机或有限状态语法从左到右的描写方式容易让人以为自然语言中存在单位线性组合,即单位是从左到右一个一个相加的。Chomsky 在 the man comes 上面加封闭圈形成 the old man comes,说明 Chomsky 可能是把 the man comes 理解成单位线性组合,即认为 the 的后继是 man,man 的后继是 comes,于是 old 插入以后成为 the 的后继:

([the old] man) comes

但实际上的线性层次是:

(the [old man]) comes

也就是说,尽管 the man comes 可以是有限状态语法生成的句子,但 the old man comes 已经不是有限状态语法生成的句子。

“父亲的父亲的父亲”可以是有限状态语法生成的句子,但也是成分线性组合:

<{([父亲]的)父亲}的>父亲

这个实例通常不会被理解成单位线性组合,因为链条中“的”和后面的“父亲”很容易看出是没有组合关系的,但在 the old man comes 中,the old 是一个可能的组合,old man 也是一个可能的组合,man comes 也是一个可能的组合,所以容易被理解成单位线性组合。事实上这只是一种错觉。

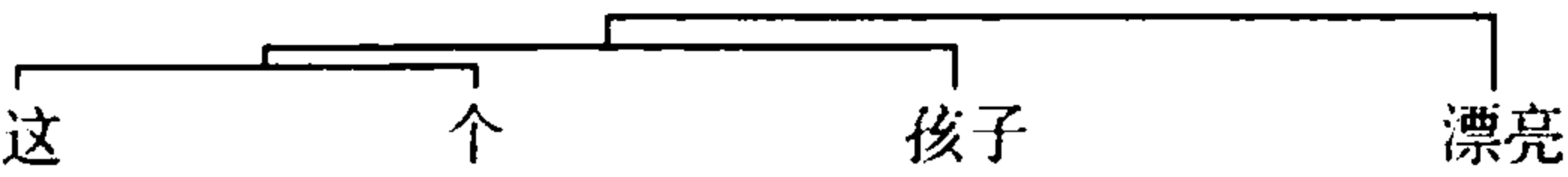
再考虑汉语的一个例子:

这个孩子漂亮
这张桌子漂亮

如果选择了量词“个”，往下就只能选择可以和量词“个”搭配的“孩子、房子、环境、窗户”等语符。一旦选择了量词“张”，往下就只能选择可以和量词“张”搭配的“桌子、画、相片”等。表面看起来这种语法选择限制似乎只是和量词这个状态有关系，但生成的组合序列仍然是有层次的，只不过都是可以用有限状态语法描写的句子。

([这个]孩子)漂亮
([这张]桌子)漂亮

即“孩子”并不是和前面的“个”直接组合，而是和“这个”组合，“漂亮”不是和前面的“孩子”组合，而是和“这个孩子”组合。这说明语法选择限制和层次并不是同一个问题，“个”和“张”限制了后面名词的不同选择，但和后面名词并不在同一个组合层次上，或者说“个”或“张”并不直接控制(immediate c-command)后面的名词。如图：



以上过程说明自然语言语符的连接本质上是成分线性过程而不是单位线性过程。根据这种性质，线性序列中的语类预测就不应该遵循单位线性过程而应该遵循成分线性过程。再考虑下面汉语的实例：

写完小说

前面说过，Markov 过程中下一个时间状态只和当前时间状态有

关系,或者说当前时间状态只和上一个时间状态有关系,这里的关键问题仍然涉及我们怎样理解 Markov 过程中的状态。如果把状态理解成单个语符,只承认单位线性过程,在不及物动词状态“完”后面(下一个状态)不可能期望名词出现。如果把状态理解成直接成分“写完”,就承认了成分线性过程,“写完”被规约成及物的 V,后面可以期望名词的出现。

Markov 链本身是可以作单位线性和成分线性两种理解的,但对于自然语言来说,如果是有限状态语法的句子,必须把从左边开始的每一次语类规约理解成状态 $q(1)$ 。如果这样做,也等于承认自然语言不存在单位线性组合。也就是说,自然语言中的直接成分问题是不可能绕开的,同时语类规约问题也不能绕开。

Chomsky 尽管没有正面论述线性问题,但他的有限状态语法的一般规则类型是:

$$A \rightarrow Aa$$

$$A \rightarrow a$$

即有限状态语法的基本规则是改写语字符串中左边的非终端语符,这样改写的结构最终是有层次的。比如,把第一个规则使用 4 次,就有:

$$\langle \{ ([Aa]a)a \} a \rangle$$

再把第二个规则使用 1 次,就有:

$$\langle \{ ([aa]a)a \} a \rangle$$

以上证明了有限状态语法或左线性语法是成分线性的。根据同

样的理由可以证明右线性语法也是成分线性的。在后面我们将论证,由于短语结构语法本质上是由左线性文法和右线性文法构成的,所以短语结构语法最终也是成分线性的。

其实有限状态自动机所能识别的线性序列也是成分线性的,只不过自动机理论中没有从这个角度分析问题。让我们从语言线性性质的角度来分析通讯理论中应用很广的奇偶校验器,这是一种典型的有限状态机,它接受的输入串是包含有 0 和 1 的符号串,当一个符号串有偶数个 1 时,输出 0,比如对于下面的符号串都要输出 0:

11
101
11011
11011000101
.....

当奇偶校验器接受的输入串包含奇数个 1 时,输出 1,比如对于下面的符号串都要输出 1:

10
10101
1101
001
11010101
.....

设计奇偶校验器只需要要两个状态就够了,即奇数状态 O(odd numbers)和偶数状态 E(even numbers)。于是有限状态图形可以描绘如下(见图 4.1.1):

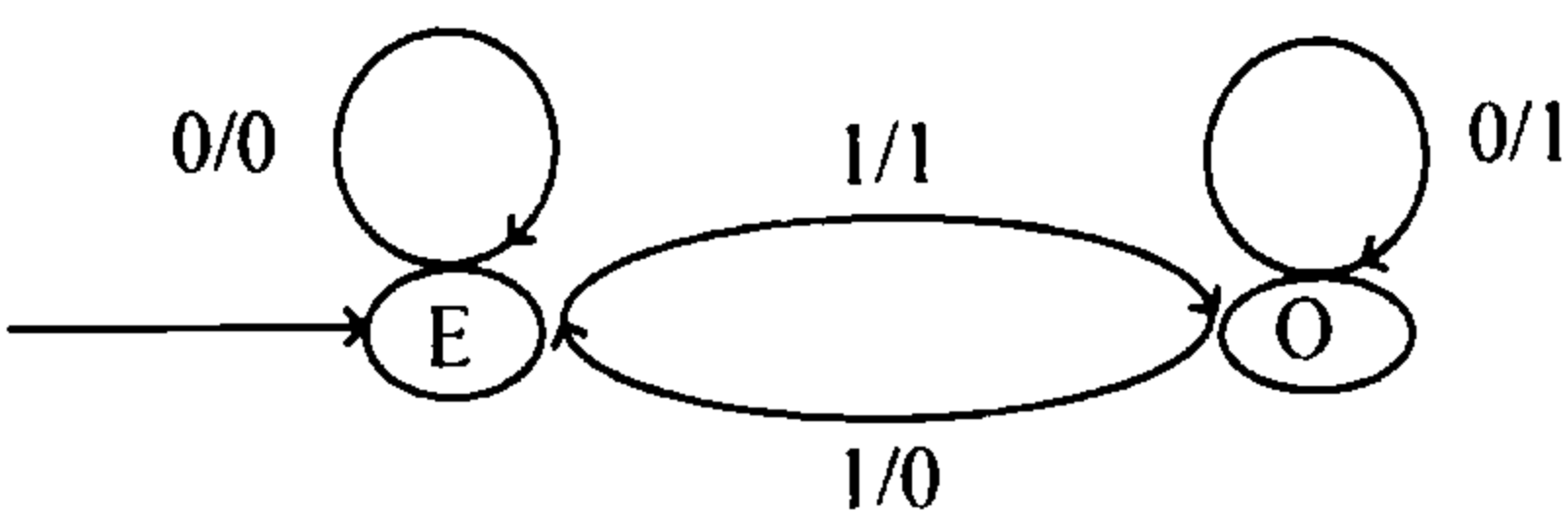


图 4.1.1 奇偶校验器有限状态图

以上斜线左边的数字表示输入的数字，斜线右边的数字表示输出的数字。这个状态图的含义是，一共只有两个状态 E 和 O，如果在状态 E 输入的是 0(在斜线左边)，输出就是 0(斜线右边)，状态停留在偶数状态 E，形成一个回路。如果在状态 E 输入 1，输出就是 1，状态 E 进入状态 O。如果在状态 O 输入的是 0，输出就是 1，状态停留在 O，形成一个回路。如果在状态 O 输入的是 1，输出就是 0，状态 O 回到到状态 E。显然，在这样的有限状态自动机中，可以识别任何包含 n 个 1 和 m 个 0 的数字串，可以判定这样的数字串是奇数个 1 还是偶数个 1。比如对于识别 0111001 这样的串，我们可以列出各个状态的机器输入和输出结果：

输入串	0	1	1	1	0	0	1
状态	E	O	E	O	O	O	E
输出串	0	1	0	1	1	1	0

由于最后达到的是 E 状态，判定该数字串是偶数。

现在的问题是，0111001 到底是单位线性的还是成分线性的？这仍然取决于我们怎样理解或设计当前状态。上面的状态是 E 和 O，实际上这种状态已经包含了成分线性的观念，我们可以通过树形图和成分标注证明这一点：

【 $E/O \mid O \langle O \{ E(O[E0]1)1 \} 1 \rangle 0 \mid 0/1$ 】

这里标注方式是和自然语言中给直接成分标注语类的方式是一致的，这里的状态 E 和 O 相当于语类，所以以上对 0111001 的识

别本质上是以成分线性方式进行的,其树形图(见图 4.1.2)为:

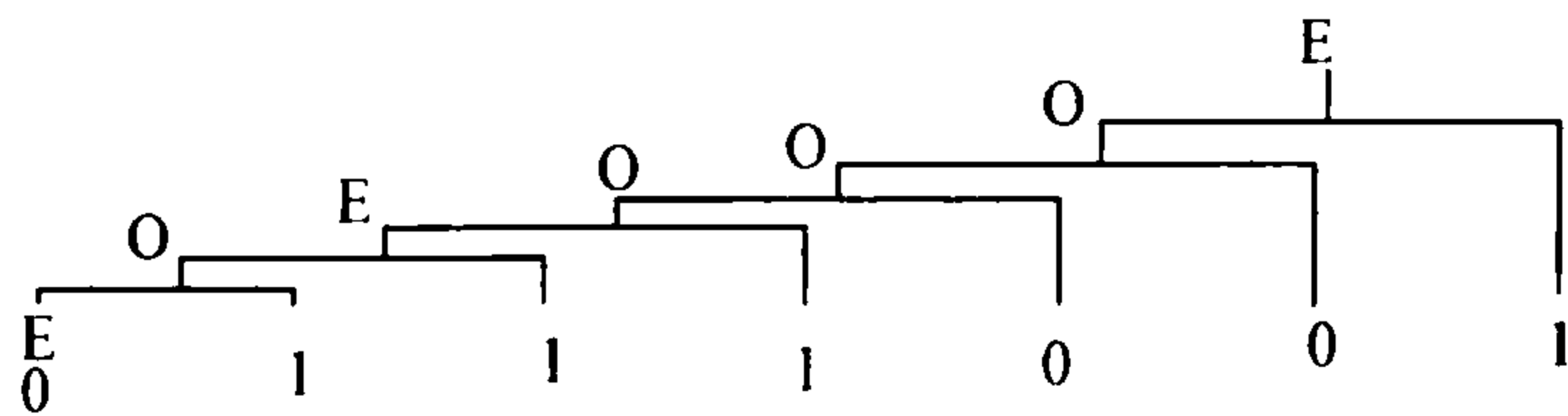


图 4.1.2 树形图

这当然满足有限状态的条件。

自然语言的成分线性关系本质上是直接成分线性关系。为了更进一步理解直接成分和有限状态自动机中状态的关系,我们再以上面的奇偶校验器为例来分析一下有限状态机的本质。一般地说,一台有限状态机 M 由六个部分组成:

- 状态集合 Q
- 有限的输入符号集合 S
- 有限的输出符号集合 R
- 状态转换函数 f
- 状态输出函数 g

奇偶校验器的输入符号集合是 $\{0,1\}$,输出符号集合也是 $\{0,1\}$,状态集合是 $\{O,E\}$,状态转换函数是:

$$f(E,0)=E,f(E,1)=O,f(O,0)=O,f(O,1)=E$$

状态输出函数是:

$$g(E,0)=0,g(E,1)=1,g(O,0)=1,g(O,1)=0$$

其中状态转换函数的性质是理解自然语言直接成分的关键。自然语言中的词或语符相当于状态输入符号。在“读书辛苦”这样

的组合中,当“读”后面输入“书”时,就进入了一个状态,这个状态并不是“书”而是“读书”,因为再下一步并不是“书”选择“辛苦”,而是“读书”选择“辛苦”,可见状态就是直接成分。从自然语言理解的角度看,有限状态自动机的状态转换函数就是从前到后不断规约直接成分的函数。在自然语言中,直接成分的实例是无限的,但直接成分的种类是有限的。

4.2 替换的有向性与有限状态语法的左线性特征

Chomsky(1957)的实例“the man comes”为什么会被误解为单位线性? 这需要进一步弄清替换的性质。考虑下面两种不同的情况:

[山东的]苹果

[山东烟台]苹果

第1个实例中,在没有“山东”的情况下,“的”并不跟“苹果”组合。第2个实例中,在没有“山东”的情况下,“烟台”也可以跟“苹果”组合,所以“山东烟台苹果”容易被误解成单位线性组合。Chomsky(1957)所用的“the man comes”之所以可能被理解为单位线性的实例,应该属于这种原因。

自然语言符号序列不存在单位线性组合,从根本上说是由扩展替换和叠加的性质决定的。下面来考虑扩展替换的各种情况。对于线性连接形成的 XY ,如果扩展替换左边的成分 X 就形成左线性组合 $[XZ]Y$,如果扩展替换右边 Y 就形成右线性组合 $X[ZY]$ 。一般的情况是:

	扩展替换方式	线性性质	原形实例	扩展替换实例
XY	[XZ]Y	左线性	写书	[写过]书
XY	[ZX]Y	左线性	房子多	[小房子]多
XY	X[Z Y]	右线性	写书	写[新书]
XY	X[Y Z]	右线性	张三在	张三[在家]

对于以上方括号中的两个符号,又可以进行扩展替换,替换的方式也是以上四种,不是左线性就是右线性,所得到的扩展序列都是有层次的。反复使用左扩展替换可以得到很长的左线性组合,前面已经有很多实例。反复使用右扩展替换也可以得到很长的右线性组合,下面是右端的“房子”不断被扩展替换,形成一个比较长的右线性组合:

画<木头{房子}>
画<木头{小(房子)}>
画<木头{小(红[房子])}>

以上只是 XY 中的某个成分不断被扩展,所以是单向线性的。还可以同时考虑两个成分同时扩展:

	扩展替换方式	线性性质	原形实例	扩展替换实例
XY	[XZ][WY]	双向线性	写书	[写过][英语书]
XY	[ZX][YW]	双向线性	写书	[常写][书摘要]
XY	[ZX][WY]	双向线性	房子多	[小房子][很多]
XY	[ZX][YW]	双向线性	房子多	[小房子][多一间]

以上每一种扩展替换都必然是有层次的,是直接成分的线性关系。可以看出,有限状态语法是替换扩展的特殊情况,即左线性扩展,因此有限状态语法所生成的序列也是有层次的,单位线性过程不可能解释由左扩展替换生成的序列。

从本质上看,扩展替换无非是左向扩展替换和右向扩展替

换,双向扩展无非是同时进行左向扩展和右向扩展,因此,符号的线性组合最终都是不断使用左向扩展和右向性规则生成的。

前面说有限状态语法生成的组合实际上是有层次的。下面进一步来分析有限状态语法和扩展替换的关系。观察下面实例:

	替换方式
[苹果]甜	
[[[烟台]苹果]]甜	苹果→烟台苹果
[[[山东]烟台]苹果]甜	烟台→山东烟台
[[[[中国]山东]烟台]苹果]甜	山东→中国山东

以上的替换特点是不断扩展替换最左边的语符,就是上面提到的左扩展替换,左扩展替换生成的结果正好是有限状态语法能够描写的语言,即左扩展替换生成的组合就是有限状态语法生成的组合。因此,从扩展替换的角度出发,我们再次看到有限状态语法只是扩展替换所生成的序列的特殊情况。由于扩展替换生成的序列是有层次的,有限状态语法也是有层次的,和单位线性组合并不相同。

以上讨论的是两个成分的扩展情况。在汉语中,有的动词后面需要两个或两个以上成分直接组合:

请小邵来

这也不能证明存在单位线性组合,因为以上的“小邵”和“来”都是成分,可以被更长的单位替换:

请[一位帮手][迅速过来]

可以说“请小邵来”是成分线性组合的特殊情况。

以上扩展替换也都和叠加是等价的,左扩展替换等于是在直接成分的右边进行右叠加。比如:

中国+山东→[中国山东]

[中国山东]+烟台→([中国山东]烟台)

([中国山东]烟台)+苹果→<([中国山东]烟台)苹果>

<([中国山东]烟台)苹果>+甜→<([中国山东]烟台)苹果>甜

右扩展替换等于是在直接成分的左边进行左叠加。比如：

红+[房子]→(红[房子])

小+(红[房子])→{小(红[房子])}

木头+{小(红[房子])}→<木头{小(红[房子])}>

画+<木头{小(红[房子])}>→画<木头{小(红[房子])}>

所以叠加形成的线性序列也都是有层次的。

单位线性组合在人工语言中是可以构造出来的。在自然语言的纯语音层面或能指层面可能存在单位线性组合。考虑下面的变调模式：

小 3+雨 3+伞 3→小 3+[雨 2 伞 3]→小 3[雨 2 伞 3]

数字表示第几调。这是有层次的变调，“雨”和“伞”是直接成分关系，“雨”变成 2 调，“小”再遇到 2 调后就不再变调了。根据我们的初步调查以及其他学者的观察，还有一种情况：

小 3+雨 3+伞 3→小 2 雨 2 伞 3

即前两个字都变调，和有层次的变调是矛盾的。可以理解成是左线性变调，即第一个字遇到上声后变成阳平，第 2 个字遇到上声后再变阳平，形成左线性变调。于是上声字序列存在层次变调和左线性变调两种情况。

之所以存在左线性变调这种情况,可能和能指层面的性质有关系,能指层面可能不一定具有扩展替换的性质。

4.3 有限状态语法的根本属性

以上讨论说明有限状态语法是扩展替换的一种特殊情况,正是因为它是扩展替换的特殊情况,有限状态语法不可能充分描写包括右线性替换在内的自然语言符号序列。Chomsky (1957)提出了短语结构语法,并证明了有限状态语法是短语结构语法的一种特殊情况。

有限状态语法必须假定符号序列从左到右的连续性。Chomsky(1957)认为,如果能够证明自然语言单位的连接过程中有非线性或不连续的现象,就能证明有限状态语法描写自然语言的限度。Chomsky(1957)正是通过自然语言有不连续性(not consecutive)的特点来论证有限状态语法的局限的。Chomsky的基本结论是正确的,但论证还存在问题,不利于对符号线性本质的理解的认识。下面我们先分析Chomsky的论证过程,然后参考上面讨论的替换的性质给出一个更简单的论证方式,更进一步说明有限状态语法的实质。

Chomsky(1957,3.2)通过形式语言和英语的类比讨论了有限状态语法的描写限度。Chomsky(1957,p21)给出了三种形式语言:

- (10) (i) ab, aabb, aaabbb.....
 (ii) aa, bb, abba, baab, aaaa, bbbb, aabbaa, abbbba.....
 (iii) aa, bb, abab, baba, aaaa, bbbb, aabaab, abbabb.....

这三种模式都不能用有限状态语法生成。比如语言(10)(i),生成规则应该是:

$$A \rightarrow aAb$$
$$A \rightarrow ab$$

从替换的角度看,是不断替换中间的成分。这显然不属于有限状态语法,前面已经提到有限状态语法的基本改写规则类型是 $A \rightarrow Aa$,是不断替换左边的成分。由于英语中存在(10i)这种句子,因此有限状态语法不能描写英语中的某些句子。Chomsky 的论证过程是先假设有下面三个句子:

(11)(i) if S1, then S2.

(ii) Either S3, or S4.

(iii) The man who said that S5, is arriving today.

在上述例子中,Chomsky 认为逗号两边的部分都有依存性(dependency),符合有限状态语法^①。Chomsky 认为把上面的 S1 用(ii)来代替,可以有:

if ,either S3, or S4, then S2.

如果再用(iii)来替代 S3,有:

(12) if ,either (11iii), or S4, then S2.

即:

if ,either the man who said that S5, is arriving today, or S4, then S2.

^① 实际上 Chomsky 在这里所列举的三个例子不完全是有限状态语法所能处理的。因为 if 和后面句子中的第一个词就没有搭配关系。

Chomsky 想用这个过程来证明英语中可以找到这样的序列：

$$a + S1 + b$$

并且还有：

$$S1 = c + S2 + d$$

进行代入就有：

$$a + c + S2 + d + b$$

从而得到镜像句子。Chomsky 认为，因为是插入，所以破坏了连续性，因此有限状态不能描写这样的句子。Chomsky 的基本思路是合理的，但列举的例子不是很合理，不便于理解有限状态语法的实质。从我们前面讨论的扩展替换的性质看，实际上“if S1, then S2”一开始就不能用有限状态语法描写，不符合有限状态语法的条件，因为这不是一个左线性扩展替换的序列。比如说：

if I were you, then I would go.

在符号层面，“if I”是不能组合的，左线性替换不可能生成出上面的句子。

其实，只要从替换的性质入手，有限状态语法的描写限度可以从自然语言中很容易得到证明。考虑下面的英语句子：

Anyone who says that is lying

其替换层次是：

([Anyone][who says that])(is lying)

anyone 和 who 没有关系,也就是说,根据有限状态语法从左到右的选择规则, anyone 不可能选择 who。更形式化地看,有限状态语法显然不能说明下面英语句子的生成过程:

Anyone	who	says	that	is	lying
W0	W1	W2	W3	W4	W5

在这里, anyone(W0)和 who(W1)没有关系, that(W3)和 is(W4)没有关系。汉语中也面临这样的问题:

我的哥哥很喜欢看电影

“我的哥哥”和“很”之间没有组合关系,“我的哥哥”是和“很喜欢看电影”组合的。在每一种语言中,都存在大量不能用有限状态语法生成的句子。

再来看 Chomsky 举的例子 the old man comes。Chomsky 认为这是一个加封闭圈后形成的例子,有限状态语法可以识别。事实上,从扩展替换形成的层次看,这个例子也不能被有限状态语法识别:

(the man)(comes)
→([the][old man])(comes)

说 the 和 old 有组合关系比较勉强。“old man”并不是对 the man 的左项“the”的扩展替换,而是对 the man 的右项“man”的

扩展,所以整个扩展过程并不是左线性扩展。类似地,如果我们在 the 后面加入一个 very,把它变成:

The	very	old	man	comes
W0	W1	W2	W3	W4

在这里,The (W0)和 very(W1)也是没有关系的,所以这样的句子也不是有限状态语法能够生成的。这里的扩展替换层次是:

$\langle_{NP} the (\langle_{NP} [_{AP} very old] [man]) \rangle \rangle \langle comes \rangle$

我们认为有限状态语法或左线性语法只能是不不断替换线性序列最左边的符号的过程,左线性序列要加封闭圈,只能加在最左边的语符上。凡是把封闭圈加在序列中间,都会形成非左线性序列。

4.4 替换、叠加与短语结构语法

Chomsky(1957)提出了另一个形式化生成模型:短语结构语法。该模型和自动机理论成为计算机科学理论中的重要内容。Chomsky 认为短语结构语法是建立在直接成分基础上的。在我们看来,一般情况是如此,但也有些区别。如果我们把叠加型直接成分和替换型直接成分区别开,可以看出短语结构语法是建立在替换型直接成分基础上的,是对替换型直接成分理论的形式化,并且在替换型直接成分的理论中注入了“生成”的思想。

在转换生成语法中,短语结构语法后来被 XP 语杠理论代替,但短语结构语法基本的替换和扩展仍然是语杠理论的基础。由于短语结构语法是直接成分理论的形式化,短语结构语法至

今还在很多其他语言模型中使用,尤其是在计算机语言理论中。下面我们来讨论短语结构语法和替换、叠加的关系,进一步认识线性组合的性质。

4.4.1 短语结构语法($[\Sigma, F]$)的基本结构

让我们来比较 Harris 和 Chomsky 面对动词短语所给出的不同描写方式:

$VP = V + NP$ (Harris, 1946)

$VP \rightarrow V + NP$ (Chomsky, 1957)

Harris 的模式,用传统语法的术语解释,就是一个动词加一个名词短语等于一个动词短语。仔细分析 Harris 等号两端的全部用例,等号表达的就是分布相等,或语类相等。这必然要追问直接成分的功能和结构的功能是否相等,这是一个相当复杂的问题。Chomsky 的 $VP \rightarrow \text{Verb} + NP$ 一类规则称为改写规则或重写规则,用来说明直接成分和结构的关系,不必受语类功能是否相同的限制。具体地说,Chomsky 的做法有两个方面的特点:

1. 改写规则把生成性彻底形式化了。
2. 改写规则是一个有序步骤过程,指明了句子生成的有限规则,所以改写规则是从 $S \rightarrow V + NP$ 开始,到终端语符列结束。

当然 Harris 的语类等价式也暗示了生成的思想,只是没有明确展开。下面来具体分析 Chomsky(1957)的一组短语结构规则:

- (1) $\text{Sentence} \rightarrow NP + VP$
- (2) $NP \rightarrow T + N$
- (3) $VP \rightarrow \text{Verb} + NP$

- (4) $T \rightarrow \text{the}$
- (5) $N \rightarrow \{\text{man, ball, ...}\}$
- (6) $\text{Verb} \rightarrow \{\text{hit, took, ...}\}$

这些规则的基本格式是 $X \rightarrow YZ$ 。左边是一个范畴符号,右边是一个以上的范畴符号或终端符号,由具体词项构成的符号序列是终端符号序列。抽象范畴符号由大写字母表示,终端语符由小写字母表示。使用以上规则,按照一定的推导式,就可以生成下面的终端符号序列:

the man hit the ball

按照上述改写规则,每次改写一个语符,the man hit the ball 的派生过程是这样的(按照规则的编号进行):

- (1) NP+VP
- (2) T+N+VP
- (3) T+N+Verb+NP
- (4) the+N+Verb+NP
- (5) the+man+Verb+NP
- (6) the+man+hit+NP
- (7) the+man+hit+T+N
- (8) the+man+hit+the+N
- (9) the+man+hit+the+ball

这个过程叫句子的派生(derivation),它是建立在短语结构规则基础上的派生。这个生成过程在后来的转换生成语法中也可以用括号标记表述:

$\{_{NP}[_T \text{ the}] [_N \text{ man}]\} \{_{VP}(V \text{ hit}) (_{NP}[_T \text{ the}] [_N \text{ ball}])\}$

这个生成过程也可以用树形图(diagram)表示(见图 4.4.1.1):

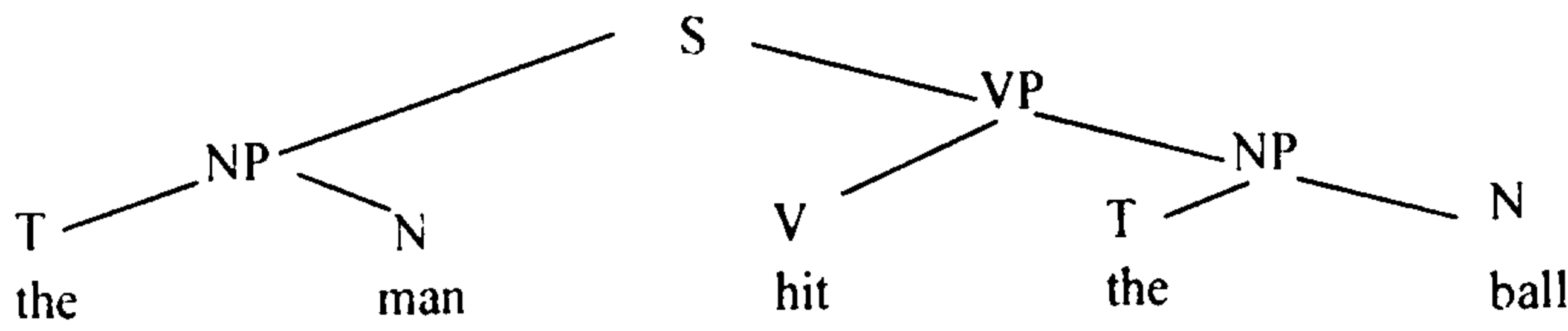


图 4.4.1.1 生成过程树形图

树形图保留了派生过程中确定短语结构(成分分析)的基本信息,更为直观,能够回溯到节点都是成分(constituent),比如 hit the ball 可以回溯到节点,是成分,man hit 不能回溯到节点,不是成分。同一节点下的两个分支构成直接成分,比如 NP(the man)和 VP(hit the ball)构成 S 的直接成分。

Sentence→NP+VP 这种描写必须以下面的信息为前提:

[the man] [hit the ball]

怎么知道不是下面的分析:

[the man hit][the ball]

Chomsky(1957)没有交代,实际上是接受了 Bloomfield、Harris、wells 的直接成分理论中对句子结构的处理方式,即在 NVN 这样的句子中,句子的直接成分在传统语法中的主语和谓语,谓语的直接成分是动词和宾语。生成语法在后来的发展中开始用 IP 这样一些概念,参看后面的讨论。

派生过程和树形图相比,树形图包含的信息比派生过程略少,因为运用规则的次序并没有反映在树形图中,所以上面的句子的树形图是唯一的,但可以有不同的派生过程。

如果两个派生可以得到同样的树形图,这两个派生就是等值的(equivalent)。对于某个特定的句子,一个句法可以允许我们构造出多个不等值的派生过程,就说这是一种结构同音(constructional homonymity)。换句话说,如果一个句子允许有两个树形图,就是结构同音,是有歧义的。

4.4.2 短语结构语法的本质

前面讨论过,从线性原则看,有限状态语法是从左到右的过程。Chomsky(1957)认为短语结构语法是从上到下的过程,需要“瞻前顾后”,需要记忆。这只是 Chomsky 的一个比喻的说法,是对元语言的一种说明。从线性原则看,短语结构语法的成分连接可以是左到右的,也可以是从右到左,也可以是两种情况的结合。

Chomsky 认为短语结构规则比有限状态语法的生成力要强。Chomsky(1957)提到的一个形式语言实例很能够说明问题:

ab, aabb, aaabbb.....

这里的句子都是由 n 个 a 和 n 个 b 组成,是无限的,生成这些句子的模式是:

$$Z \rightarrow ab$$

$$Z \rightarrow aZb$$

a 和 b 有直接组合关系,但 a 和 a 没有直接组合关系, b 和 b 也没有直接组合关系,所以有限状态语法不能生成上述全部句子。从改写规则可以看出,插入 a 和 b 之间的 Z 隔断了这种联系,所以这样的句子只能用短语结构语法生成。

在实际语言中,能用短语结构语法生成但不能用有限状态语法生成的句子也很多。比如在汉语中,“他很高”中的“他”和“很”由于不能直接组合,所以不能用有限状态语法生成。而在短语结构语法中,只要有规则:

$$S \rightarrow NP + VP$$
$$VP \rightarrow Adv. + Adj.$$

就可以生成“他很高”这样的句子。

短语结构规则的基本形式是: $X \rightarrow YZ$ 。改写规则箭头左边的语符只有一个,右边有两个。短语结构规则实际上接受了人类在生成句子时要涉及功能范畴(词类)的归约或语类归约这种思想,因此也涉及直接成分的思想。有限状态语法只考虑左线性层次。由于“他很高”这样一批句子不能通过有限状态语法生成,而需要借助短语结构语法,又由于 Chomsky 的短语结构语法建立在直接成分理论或层次的基础上,由此可以进一步证明单位在线性序列中的组合层次是自然语言的一种初始性质。

Chomsky(1959, 4.2)提出了下面的定理:

每一种有限状态语言都是一种终端语言,但是有些终端语言却不是有限状态语言。

Chomsky(1959)还进一步讨论了形式语言的类型。

一般认为,而且我们前面也暂时认定,短语结构语法是建立在直接成分理论基础上的。其实问题还不是这么简单,这取决于我们怎样理解直接成分。前面讨论过,直接成分有两种情况:

我/读过《阿 Q 正传》

《阿 Q 正传》/我读过

如果我们仔细分析 Chomsky(1957)短语结构规则的细节和所涉及的实例,只有“我读过《阿 Q 正传》”这种语序可以直接通过短语结构规则生成,“《阿 Q 正传》/我读过”这种语序的生成则需要进一步考虑转换规则。把哪些语序作为短语结构规则生成的结果或作为基础部分生成的结果,一直是一个有争议的问题。有很多学者试图找出理论上的界定标准。我们后面论证转换的必要性时会谈到,广义结构语法等模型试图把全部的话语作为短语结构生成的结果。我们认为,上面两种语序都可以通过基式和叠加产生,但只有前一个实例可以通过基式和扩展替换生成。一般地说,Chomsky 的短语结构语法的所有规则,都可以通过左扩展替换、右扩展替换、双向扩展替换和等长替换得到解释。比如,前面提到的 Chomsky(1957)短语结构语法的各个改写规则,都是一种扩展替换的结果:

(1) Sentence $\rightarrow$ NP+VP	扩展替换,可继续双项扩展替换 (NP 和 VP 都可以扩展)
(2) NP $\rightarrow$ T+N	扩展替换
(3) VP $\rightarrow$ Verb+NP	扩展替换,可继续右项扩展(NP 可以扩展)
(4) T $\rightarrow$ the	等长替换,词汇插入
(5) N $\rightarrow$ {man,ball...}	等长替换,词汇插入
(6) Verb $\rightarrow$ {hit,took...}	等长替换,词汇插入

所以从根本上说,短语结构规则是连接和替换的形式化,而不是连接和叠加的形式化,因此所生成的序列同替换所生成的序列一样,不是左线性的,就是右线性的。比如,替换中的右扩展体现在短语结构语法的以下规则中:

(2) $NP \rightarrow T + N$

(3) $VP \rightarrow Verb + NP$

可生成右线性序列:

$Verb + [T + N]$ (hit[the ball])

如果 VP 规则中的 NP 为零,左扩展体现在短语结构语法的以下规则中:

$S \rightarrow NP + VP$

$VP \rightarrow Verb$

$N \rightarrow T + N$

可以生成左线性序列:

$[T + N] + Verb$ ([the man]comes)

由于短语结构语法实际上是连接和替换的形式化,并不是连接和叠加的形式化,所以类似“《阿 Q 正传》我读过”这样的组合语序,在 Chomsky 几个阶段的理论模式中,都是放在转换中来处理的,但并没有明确给出理由。我们认为,一般地说,对于如下类型的改写规则模式:

$X \rightarrow YZ$

.....

$A \rightarrow a$

$B \rightarrow a$

.....

由于最终要改写到两个单位组合,而两个单位的组合相当于我们前面讨论的初始连接,所以短语结构规则都是连接和替换的形式化,不是连接和叠加的形式化。短语结构规则并没有描写由叠加生成的有层次的组合。这些由叠加生成的、可以用直接成分来分析的组合,转换生成语法都是放到转换中处理的。从中国学者所作的层次分析中,我们可以分析出两种情况,一种层次分析是可以通过替换解释的,比如“我/读过《阿 Q 正传》”,一种层次只能通过叠加解释,比如“《阿 Q 正传》/我读过”。

有人认为自然语言的描写可以不要转换规则,只需要对短语结构规则加以扩充就够了,如广义短语结构语法(Generalized Phrase Structure Grammar)、接口语法(Interface Grammar)、Joan Bresnan 的词汇-功能语法(Lexical-Functional Grammar)、关系语法(Relational Grammar)、弧对语法(Arc Pair Grammar)等。现在我们从叠加的角度来分析这个问题。由于线性结构有一部分是通过叠加生成的,可以考虑把这部分线性组合用短语结构规则处理,这时候的短语结构规则是扩展的短语结构规则:

[我][读过《阿 Q 正传》]

[《阿 Q 正传》][我读过]

我们可以用 Gazdar 和 Pullum(1985)的广义短语结构语中的两个短语结构规则:

$$S \rightarrow NP VP$$

$$S \rightarrow NP S/NP$$

得到:

[_{NP}我][_{VP}读过《阿 Q 正传》]

[_{NP}《阿 Q 正传》][_S我读过]/NP

后一个规则中的“S/NP”即所谓斜线范畴(slash category),表示句子中缺一个名词短语,是为了扩展短语结构规则而引入的。这样一来原来由转换规则处理的叠加线性组合就由广义短语结构规则处理了。现在我们可以提出两个问题:

1. 在带有斜线范畴的规则中,或类似斜线范畴的规则中,是否已经暗含了转换规则的性质?斜线后的名词短语很像Chomsky后来引入的空范畴(empty category)。Chomsky的空范畴是和移位- α (Move- α)密切相关的概念,而移位- α 就是对转换规则的高度概括。这暗示所有的转换规则根本上就是在解释成分的移位以及移位后所留下的痕迹(trace)。上述包含斜线范畴的规则与移位- α 有相似之处,这可能暗示广义短语结构语法中的有些规则并没有真正越出转换规则的实质。

2. 即使不承认空语类,对那些不是通过叠加形成,而是通过替换加转换形成的序列,是否都可以用广义结构语法解释?比如这样的句子:“小邵将汪锋的军。”

我们将在后面讨论深层结构的时候再回答这里的问题。

关系语法取消转换的一个重要理由是世界上有些语言的被动句和主动句的差别不是位移的结果,而是标记的变化,即被动句和主动句有相同的语序。问题在于,有些语言没有主动句和被动句的语序差别并不等于人类其他语言没有这种差别。更重要的是,转换作为一种生成句子的手段,并不只是解决被动句的生成问题。只要人类语言存在必须用转换来生成的序列,转换就是必要的。

4.5 有限状态语法和语类规约模型

根据 Chomsky(1957)的理论,短语结构语法的理解机制是从上到下的,有限状态语法是从左到右的。从句子生成的角度来看,Chomsky 的这种认识应该没有问题,而且我们可以进一步把有限状态语法看成是从上到下的一种特殊情况。尽管有限状态语法或左线性语法的生成能力有限,但从理解句子的角度看,由于符号的能指序列是按照先后顺序接受的,而有限状态语法的左线性特征就是按照先后顺序描写,所以有限状态语法仍然是一个最贴近理解过程的模型,也最贴近语符在时间上先后输入的过程,即语符是从左到右一个一个输入的。从左到右理解线性组合仍然是一个符合能指顺序的过程。如果我们能够在左线性语法的基础上作一些扩展,或许能够对自然语言的理解模型有更好的解释。下面我们尝试从有限状态语法的观点出发,给出一种语类规约模型来识别句子的语类。识别语类是识别句子的必要步骤,正确的理解不应该违背语类规约。一个句子的语类规约可以有多种可能,在此基础上再引入语义、语音、语用、认知等标准,就可能最终理解和生成合法的句子。

纯粹的有限状态语法从左到右的左线性语法特征决定了不能识别有些句子,比如前面提到的典型例子:

房子[很宽敞]

这是一个右线性句子,单位是从右往左先组合,所以如果从左到右,“房子”遇到“很”不能完成组合。我们需要考虑在面临这样的问题时怎样继续从前往后理解句子。前面讨论过,在连接的基础上,叠加可以描写全部线性句子,由于叠加是有层次的,最终都是直接成分的叠加,根据前面讨论过的语类推导性质,在语

符或词的基础上,通过直接成分的语类可以推导结构的语类,结构作为更大的结构的直接成分,又可以推导高层结构的语类,所以线性理解过程本质上是语类规约的问题。下面考虑三个语符的情况。一种是左线性理解方式,语类规约过程是:

$$\begin{aligned} & [D \text{ 金}][N \text{ 镯子}][A \text{ 贵}] \\ & \rightarrow (NP[D \text{ 金}][N \text{ 镯子}])(A \text{ 贵}) \\ & \rightarrow \langle_S(NP[D \text{ 金}][N \text{ 镯子}])(A \text{ 贵}) \rangle \end{aligned}$$

以上语类规约可以抽象的概括为:

$$[D][N][A] \rightarrow [NP][A] \rightarrow S$$

一种是右线性理解方式:

$$\begin{aligned} & [N \text{ 手表}][Ad \text{ 很}][A \text{ 贵}] \\ & \rightarrow (N \text{ 手表})(AP[Ad \text{ 很}][A \text{ 贵}]) \\ & \rightarrow \langle_S(N \text{ 手表})(AP[Ad \text{ 很}][A \text{ 贵}]) \rangle \end{aligned}$$

以上语类规约可以抽象的概括为:

$$[N][Ad][A] \rightarrow [N][AP] \rightarrow S$$

一般地说,处在两个语符之间的语符 X,有前后组合两种方式。在第一例中,中间的“镯子”和前面的“金”组合,语类规约为 NP,NP 再和后面的 A(贵)组合,语类规约为 S。第二例中,中间的“很”和后面的“贵”组合,语类规约为 AP,前面的 N(手表)再和 AP 组合,语类规约为 S。更多的语符组合,基本原则也是相同的。归纳起来,链条中的 X 存在前后组合的可能。在有限

步骤中,语类规约总是能够完成的。

以上两个实例实际上都是在用叠加,区别在于第一个实例的叠加和左线性语法是等价的,第二个实例由于“手表”跟“很”不能组合,“很”只能和后面的成分叠加,这时候“手表”有一个暂停状态,等“很”和“贵”组合成 AP 后,“手表”再和 AP 组合。再考虑下面实例:

小邵学习的书房很宽敞

如果采取从上到下的短语结构理解模式,首先需要确定句子的 NP 和 VP 的连接点,这个连接点显然不在“小邵”和“学习”之间,而在“书房”和“很”之间。怎样判定这一点,Wells(1947)提出了扩展法。我们在后面会看到,扩展法仍然要依赖从左到右的条件,因此从上到下的过程并不能提供直接从左到右的可行方案,我们需要的是从左到右不断对状态进行语类规约。

如果在叠加的基础上加上暂停状态,就能够解决从左到右的可行性问题。为了说明这个问题,我们先从前面讨论过的简单实例“苹果甜”开始。“苹果甜”可以进行最右端符号的扩展:

苹果[甜]→苹果[很甜]

“苹果”跟“很”是不会发生线性组合的。再把“苹果”的扩展结果加在一起,有:

<{([中国]山东)烟台}苹果>/<很甜>

有限状态语法或左线性语法一直可以理解到“苹果”,但在“很”前面理解受阻。如果允许有暂停状态,有限状态语法也能理解上面的句子,即在遇到“很”时,记录下前面的 NP 状态:

[_{NP}中国山东烟台苹果][_{AP}很甜]

然后从“很”开始继续往右叠加,得到“很甜”,记录下 AP,再退回到 NP 状态,使用叠加,识别出[NP]和[AP]叠加是合法序列。

更一般地说,以上实例的理解过程反映了语类序列的一种理解过程:

中国山东烟台 苹果很甜		语类归约 规则
[N][N][N][N] [Ad][A]	中国山东烟台苹果很甜	
[NP][N][N] [Ad][A]	[_{NP} 中国山东]烟台苹果很甜	N+N→NP
(NP)[N][Ad] [A]	(_{NP} [_{NP} 中国山东]烟台)苹果很甜	NP+N→NP
{NP}[Ad][A]	{ _{NP} (_{NP} [_{NP} 中国山东]烟台)苹果} _{很甜}	NP+N→NP
{NP}{AP}	{ _{NP} (_{NP} [_{NP} 中国山东]烟台)苹果}{ _{AP} 很甜}	Ad+A→AP
S	< _S { _{NP} (_{NP} [_{NP} 中国山东]烟台)苹果}{ _{AP} 很甜}>	NP+AP→S

按照语类规约,前面几个 N 都可以通过名词短语(名词)加名词仍然是名词短语这样一个语类规则识别,但当 NP 和 Ad 相遇时,不符合语类组合条件,有限状态语法失效,说明这个语类序列不完全是左线性的。这时候可以记录一个暂停状态 NP,然后可以从 Ad 开始继续往右叠加,由于 Ad 和后面的 A 可以组合,构成 AP,这时可以记录第 2 个暂停状态 AP。然后在第一个暂停状态 NP 和第 2 个暂停状态 AP 之间再使用叠加,由于 NP 和 AP 可以组合,识别完成,证明 NNNNAdA 是合法的汉语

语类组合。

以上分析说明,尽管 NNNNA_dN 不完全是左线性的,但仍然可以通过暂停状态,把整个序列分成两个左线性序列,各自进行左线性组合,最后对两个子序列进行左线性组合。以上语类序列的暂停状态 NP 和 AP,其改写规则可以描写为:

$$S \rightarrow NP \ AP$$

这实际上就是 Chomsky(1957)短语结构语法 $X \rightarrow YZ$ 的一种具体情况。从这种意义上说,右向叠加或左线性组合如果允许暂停状态,其描写能力就已经包含了上下文无关文法或短语结构语法的描写能力。

由于用叠加生成的线性组合包括了用扩展生成的线性组合,既然语类规约模型能够理解叠加生成的所有线性模型,因此也能理解所有通过扩展形成的线性模型。这一点也可以直接通过扩展的机制得到证明。一般地说,对于下面的扩展替换:

$$XY \rightarrow [X_1X_2][Y_1Y_2] \text{ (苹果甜} \rightarrow \text{[红苹果][很甜])}$$

其中 X 和 Y 都分别被替换成 X_1X_2 和 Y_1Y_2 。显然这样的替换结果不仅仅是左线性的,也不仅仅是右线性的,因为 X_2 和 Y_1Y_2 没有关系。但 X_1X_2 可以还原为 X, Y_1Y_2 可以还原为 Y,仍然可以在语类规约的基础上用左线性语法识别。一般地说,由于用替换生成的序列最终都可以分析成两个语类组成的序列,而两个语类组成的序列最终都可以用左线性语法加识别,因此由替换生成的序列最终都可以用左线性语法识别,条件是必须有暂停状态。于是我们得到下面一个结论:

如果允许暂停状态,所有通过短语结构语法理解的序列也

可以通过左线性语法或有限状态语法理解。

再回到“小邵学习的书房很宽敞”这个更复杂的实例上来。语类规约模型理解这个实例可以分成以下步骤：

$[_N \text{ 小邵 }][_V \text{ 学习 }][\text{ 的 }][_N \text{ 书房 }][_{Ad} \text{ 很 }][_A \text{ 宽敞 }]$
 $\rightarrow (_{VP} \text{ 小邵学习 })[\text{ 的 }][_N \text{ 书房 }][_{Ad} \text{ 很 }][_A \text{ 宽敞 }]$
 $\rightarrow (_{VPd} \text{ 小邵学习的 })[_N \text{ 书房 }][_{Ad} \text{ 很 }][_A \text{ 宽敞 }]$
 $\rightarrow (_{NP} \text{ 小邵学习的书房 })[_{Ad} \text{ 很 }][_A \text{ 宽敞 }]$
 $\rightarrow (_{NP} \text{ 小邵学习的书房 })(_{AP} \text{ 很宽敞 })$
 $\rightarrow \{_S \text{ 小邵学习的书房很宽敞 }\}$

以上 N(小邵)和 V(学习)可以组合成语类 VP(小邵学习), VP 接着和后面“的”可以组合成 VPd(小邵学习的), VPd 接着和后面的 N(书房)可以组合成 NP(小邵学习的书房)。到此为止都是有限状态语法或左线性语法的识别。NP(小邵学习的书房)再往下遇到“很”,不能用左线性语法识别,这时将 NP(小邵学习的书房)作为第 1 暂停状态,然后从“很”开始继续进行左线性语法的识别,Ad(很)和 A(宽敞)可以组成形容词短语 AP(很宽敞),作为第 2 暂停状态。最后第 1 暂停状态 NP 和第 2 暂停状态 AP 再组合,还原为 S,完成理解。最后整个句子的层次是:

$\text{【}_S \langle \text{NP} \{ \text{VPd} (\text{VP} [_N \text{ 小邵 }][_V \text{ 学习 }]) (\text{ 的 }) \} \{ \text{N} \text{ 书房 } \} \rangle \langle \text{AP} [_{Ad} \text{ 很 }][_A \text{ 宽敞 }]\rangle \text{】}$

这里的语类规约过程可以更一般地概括为:

$[_N][_V][_d][_N][_{Ad}][_A]$
 $(VP)[_d][_N][_{Ad}][_A]$

{VPd}[N][Ad][A]
<NP>[Ad][A]
<NP><AP>

【S】

其中 N 和 V 组合后推导出语类 VP,同时形成一个 VP 状态, VP 和右边的 d 组合后推导出语类 VPd,同时形成 VPd 状态, VPd 状态和后面的 N 组合形成语类 NP,同时进入 NP 状态, NP 和后面的 Ad 不能组合,NP 状态这时处于暂停状态。识别模型这时从 Ad 开始。Ad 和后面的 A 组合后推导出 AP 语类,同时形成 AP 状态。然后前面的 NP 状态再重新激活,和后面的 AP 组合后形成语类 S。这时全部语符已经识别完,整个识别过程结束。

现在考虑状态语类规约模型如何识别歧义。观察实例：

三个研究所的/研究员
三个/研究所的研究员

这两个实例的语类序列是：

[Q][L][N][d][N]

第一种语类规约方式：

[Q][L][N][d][N]	三个研究所的研究员	
[QP][N][d][N]	[_{QP} 三个]研究所的研究员	Q+L→QP
[NP][d][N]	(_{NP} [_{QP} 三个]研究所)的研究员	QP+N→NP

续表

[NPd][N]	{NPd(NP[QP三个]研究所)的}研究员	NP+d→NPd
[NP]	⟨NP{NPd(NP[QP三个]研究所)的}研究员⟩	NPd+N→NP

第二种语类规约方式：

[Q][L][N][d][N]	三个研究所的研究员	
[QP][N][d][N]	[QP三个]研究所的研究员	Q+L→QP
[QP][Nd][N]	[QP三个][Nd研究所的]研究员	N+d→Nd
[QP][NP]	[QP三个][NP[Nd研究所的]研究员]	Nd+N→NP
[NP]	⟨NP[QP三个](NP[Nd研究所的]研究员)⟩	QP+NP→NP

两种方式的主要差别在于进入[QP]状态后,后续的 N 可以和前面的 QP 组合成 NP,也可以和后面的 d 组合成 Nd。也就是说,在进行左线性规约的过程中,序列中任何语符都应该允许左向组合和右向组合两种可能性。如果有一种向度不能实现,就不存在歧义。这种双向可组合性也仍然可以作为左线性来处理。一般来说,对于 AXB,如果 X 可双向组合,其中一种情况包含 AX 左线性,另一种情况则是包含 XB 左线性。上面的实例也属于这种情况。

上面的 NP+d 也有人认为可以是 NP,由于 NPd 和 NP 在作定语方面有些区别,我们把 NPd 和 NP 分开。更一般地说,我们把 XPd 和 XP 分开。

再考虑一个例子：

咬主人的/狗
咬/主人的狗

两者的语素序列类是：

[V][N][d][N]

第一种规约方式：

[V][N][d][N]	咬主人的狗	
(VP)[d][N]	(_{VP} 咬主人)的狗	V+N→VP
{VPd}[N]	{ _{VPd} (_{VP} 咬主人)的}狗	VP+d→VPd
⟨NP⟩	⟨ _{NP} { _{VPd} [_{VP} 咬主人]的}狗⟩	VPd+N→NP

V+N 还可以规约成语类 NP(如“烤白薯”),“咬主人”不属于这种情况。再看第二种规约方式：

[V][N][d][N]	咬主人的狗	
[V](Nd)[N]	咬(_{Nd} 主人的)狗	N+d→Nd
[V]{NP}	咬{ _{NP} (_{Nd} 主人的)狗}	Nd+N→NP
⟨VP⟩	⟨ _{VP} 咬{ _{NP} (_{Nd} 主人的)狗}⟩	V+NP→VP

再考虑一个更复杂的实例：

傅刚知道小邵学习的书房很宽敞

“知道”后面可以是一个子句,也可以是 NP。先看第一种规约方式：

(_s 傅刚知道[_s 小邵学习的书房很宽敞])

当“知道”输入时,马上识别后面是否可以形成一个句子。事实

上后面正好可以理解成一个句子。具体过程是, N(小邵)和 V(学习)组合成 VP, VP(小邵学习)和后面的“的”组合成 VPd(小邵学习的), VPd 和后面的 N(书房)组合成 NP(小邵学习的书房), 这时 NP 不能和后面的“很”组合, 处于暂停状态, 识别重新从后面的 Ad(很)开始, Ad(很)和后面的 A(宽敞)组合成 AP(很宽敞), 这时重新激活 NP(小邵学习的书房), NP 和 AP 组合成 S(小邵学习的书房)。这时 V(知道)再和 S 组合, 形成 VP, N(傅刚)再和 VP 组合, 形成 S, 完成识别。

由于“知道”也可以和 N 组合, 所以整个句子的识别一开始也可能进入下面状态:

[_s 傅刚知道小邵][_v 学习]的书房很宽敞

这里的 S 和后面的 V(学习)不能组合(即使把“学习”识别成 N, 也不能组合), S 可以作为暂停状态, V(学习)继续和右边的“的”组合, 形成 Vd(学习的), Vd 再和后面的 N(书房)叠加, 形成 NP(学习的书房), 这时 NP 既不能和前面的 S 组合, 也不能和后面的 Ad(很)组合, 成为第 2 暂停状态。规约过程这时从 Ad(很)开始, Ad 和 A(宽敞)组合成 AP(很宽敞)。这时的第 2 暂停状态 NP(学习的书房)和 AP 组合形成 S1。第一状态 S 和 S1 不能组合, 所以整个理解过程不成功。这时需要退回到“小邵”以前, 把“傅刚知道”作为暂停状态, 如果后面能还原出一个子句, 理解就是成功的。这正是前面提到的第一种理解方式。

以上讨论的语类规约模型仅仅涉及左线性组合, 其实语类规约模型也可以是右线性组合, 只是右线性组合和自然语言语符的线性输入方向不一致, 所以在右线性组合上建立语类规约程序更复杂。不过这一事实反映了一个本质, 语类规约模型可以从序列中的任何一个语符开始向两个方向展开。由于动词在句子中处于核心位置, 一种方便的办法是从 V 开始展开语类规

约模型的理解。

现在回到线性问题的实质上来。前面我们说自然语言符号层面的线条性本质上是直接成分的线条性。现在我们看到,只要借助暂停状态,直接成分仍然可以通过有限状态语法识别,而有限状态语法本质上是左线性的,因此自然语言的线性符号序列可以得到左线性描述。

思考问题

1. 自然语言中有没有单位线性组合?
2. 根据有限状态语法识别线性句子有什么价值?

5. 线性原则的不充分性和转换的必要性

5.1 Harris 的线性化工作

我们曾经在 Hockett(1954)的基础上讨论过 IA 模式和 IP 模式的优劣(陈保亚,1999,5.7),IP 依赖直觉分析,承认通过直觉把握得到的基本单位和关系,比如词和结构关系,而把复杂片断看成是基本形式通过变化或过程得到的;IA 强调纯正的客观描写,力图不依赖直觉,只依赖观察事实,因此只承认通过对比得到的单位,即语素,并通过语素的分布而得到语素类。IA 的精神实质可以用 Hockett(1954,1.5)的一段话作为代表:

(One assumes that any utterance in a given language consists wholly of a certain number of morphemes, in a certain arrangement relative to each other.

这里的核心思想可以概括成“话语由语素配列构成”。现在我们进一步注意到,IA 模式和 Saussure 的线性思想更一致。

IA 模式有一个发展过程。Bloomfield(1933,10.4,p163)认为:

The meaningful arrangements of forms in a language constitute its *grammar*.

Bloomfield 列出了四种配列方式:

1. order(词序)

2. modulation(变调)
3. phonetic modification(变音)
4. selection of forms(形式的选择)

以上 1、2 两种方式和“话语由语素配列构成”有不一致的地方。第 3 种方式 Bloomfield 作了修正,用交替形式(alternant)来解释变音问题,交替形式的本质就是语素变体。第 4 种方式是语素类。后来 Harris(1946)又对 2 作了修正,把变调的结果作为语素处理。在 Harris(1946)的语素话语模型中,只有语素、序列类(相当于 Bloomfield 的 selection of forms)和语素、语素序列在一定位置中的替换。更准确地说,Harris 用语素和语素序列在一定位置中的替换把上述四种配列方式统一成为一种更简单的模式。

因此,IA 的典范应该是 Harris(1946)中所表述的模式。在这种模式中,只考虑语素和语素序列及其分布。这一模式蕴含了以下思想:

1.* 话语是完全是由语素配列构成的,每一段话语都可以通过对比切分出语素,每一段话语都必须能够通过对比切分出语素。换个角度说,话语只能由语素构成。

2. 由于每一段话语的形式都是一段语音,所以语素由音位构成。

于是语素组合所产生的各种变化,通过语素的切分、归并和归纳,被彻底线性化了。

5.2 转换问题的提出

句式间的转换(transformation)关系在传统语法中已经提

到。Harris 也注意到纯线性化的分布分析的不足,首先展开了关于转换的理论分析(Harris,1952)。Harris 认为 the weather is cold 和 the cold weather 有联系,或者说 the N is A 和 the A N 有联系。Harris(1957, P288)对转换的定义是:

If two or more constructions (or sequences of constructions) which contain the same n-classes (whatever else they may contain) occur with the same n-tuples of member of these classes in the same sentence environment, we say that the constructions are transforms of each other, and that each may be derived from any other of them by a particular transformation.

Harris 这里的 n-tuples 是指具体的语素,这是转换的条件。换句话说,相同的转换有相同的词项。从 Harris(1957,p288)的例子可以看出他对转换的理解:

N1 v N2 $\longleftrightarrow$ N1's Ving N2
He met us $\longleftrightarrow$ his meeting us
The foreman put the list up $\longleftrightarrow$ the foreman's putting the list up

以上箭头两端的片断都是可以进行直接成分分析的,因此,Harris 不是因为直接成分在分析结构的时候遇到了困难才引入转换的,而是因为不同句式有共同点才引入转换。这和 Chomsky(1957)转换生成语法中的转换思想有很大的区别。Harris (1957,p283)的转换的基础是:two structures having the same set of individual co-occurrences。

Harris(1957,p288)谈到了转换的两种情况。一种可以朝

两个方向转换的：

N1 v V N2	N2 v be Ven by N1
The kids broke the window	The window was broken by the kids
The detective will watch the staff	The staff will be watched by the detective

另一种只能朝一个方向转换的，Harris 称为单向转换 (one-directional transformations) 或者不可逆转换 (nonreversible transformations)：

N1 v V N2	N2 v be Ven by N1
*	The wreck was seen by the seashore

Chomsky(1957,p44)给出了转换的另一种解释，认为语法转换 T 是对带有既定成分结构的语素串的操作，并把该语素串转变成具有新的成分结构的语素串：

A grammatical transformation T operates on a given string (or, as in the case of (26), on a set of strings) with a given constituent structure and converts it into a new string with a new derived constituent structure. (Chomsky,1957,p44)

下面是被动式的转换规则(Chomsky,1957, p43)：

(34) If S1 is a grammatical sentence of the form
NP1 - Aux - V - NP2
then the corresponding string of the form
NP2 - Aux + be + en - V - by + NP1

is also a grammatical sentence

我们认为 Chomsky(1957)的 Syntactic Structure 在语言认识论上最重要的意义就在于提出了把转换规则作为生成语法的一个重要组成部分。尽管在后来的生成语法研究的很多文献中,“移位”常常代替“转换”这个术语,但转换作为生成语法的一个部分,这一基本框架并没有变化,由此可见《句法结构》中提出转换的重要意义。

也有些模型提出取消转换部分,其中最主要的原因在于,转换在自然语言理解中处理起来太复杂。20 世纪 70 年代末以来,在生成语法学派的圈子中,有人甚至认为自然语言的描写可以不要转换规则,只需要对短语结构规则加以扩充就够了,前面我们提到的广义短语结构语法(Generalized Phrase Structure Grammar)、接口语法(Interface Grammar)、词汇-功能语法(Lexical-Functional Grammar)、关系语法(Relational Grammar)、弧对语法(Arc Pair Grammar)都是这种思路。

从科学研究或学术研究的角度看,语言研究应当有自己的本体目标和认识论上的价值,这样才能认识语言的机制。自然语言的语法规则是否要包含转换部分,需要有科学的论证,至于在信息处理等具体应用中是否可以用其他办法来代替转换,是一个技术和策略问题。一般地说,在线性规则和转换规则都可以描写清楚语言事实的情况下,用线性规则而不用转换规则的模式似乎更简单。如果人类的语法真的不需要转换部分,那就必须承认线性程序可以生成全部句子。问题是先要证明线性规则是否具有充分性。下面我们要证明线性规则是不充分的,引入转换规则是必要的。

问题需要重新回到转换规则的论证上来,因为 Chomsky 尽管把转换引入语法部分,但我们发现 Chomsky 论证转换存在的理由还不能构成转换存在的必要条件。Chomsky(1957)是从简

单性问题上论证短语结构语法的不充分性(inadequate),也就是从评价程序上论证短语结构语法的不充分性。所以 Chomsky 在论证转换理由的时候本质上还是从描写的简单性入手的(Chomsky, 1957, 5.1, p34)。Chomsky(1957, 5.3, p41)在一个脚注里说得很明确,如果我们用扩展短语结构语法来描写所有的英语句子,就会失去描写的简单性:

It appears to be the case that the notions of phrase structure are quite adequate for a small part of the language and that the rest of the language can be derived by repeated application of a rather simple set of transformations to the strings given by the phrase structure grammar. If we were to attempt to extend phrase structure grammar to cover the entire language directly, we would lose the simplicity of the limited phrase structure grammar and of the transformational development. (Chomsky, 1957, 5.3, p41)

在 Chomsky 以后的论著中,转换直接被作为一个不可缺少的层面,其必要性没有再加以论述。我们认为正是 Chomsky 论证转换不是从必要条件入手,所以后来有人提出了取消转换的主张。

5.3 Chomsky 关于转换存在的理由

Chomsky(1957, p35)提到一种连接过程(process of conjunction),认为这一过程要参考派生史(history of derivation),属于转换性质。Chomsky 所说的参考派生史就是参考上下文。实例如下:

(20)(a) the scene—of the movie —was in Chicago

(b) the scene--of the play—was in Chicago

(21) the scene—of the movie and of the play was in Chicago

根据 Chomsky 的思路,以上生成过程本质上是转换:

$$(Z + X + W) - (Z + Y + W) \rightarrow Z - X + \text{and} + Y - W$$

但是,我们认为以上生成过程可以用线性程序生成。of the movie 和 of the play 都可以通过我们前面提到的线性叠加生成,这两个大的部分又可以叠加成 of the movie and of the play,然后再和 the scene 叠加就形成:

$$\{ \text{the (scene [of the movie and of the play])} \}$$

扩充也能解释这里的生成过程。用 of the movie and of the play 来扩充(20)中的 of the movie 或 of the play,就能生成(21)。更概括地说,用 X and Y 扩充 X 或 Y,就能生成 $Z - X + \text{and} + Y - W$,Chomsky 用转换来描写这里的连接过程,是一个优化性问题,并不是必要性问题。

Chomsky 之所以认为连接必须借助转换,和他的形式化方式有关系,他的短语结构规则(1957)都是单句的描写规则,于是两个单句的连接必须要用到上面的转换公式,并通过删除 W 部分形成新的句子。如果我们在 Chomsky(1957)的短语结构规则加上一条改写规则:

$$PP \rightarrow PP + PP$$

即介词短语加上介词短语还是介词短语,这时改写规则就可以

完成以上转换规则所完成的工作。

Chomsky(1957)还从歧义解释的角度论证了转换层面的存在理由。实际上这也是在论证优化性。对于下面的实例(Chomsky,1957,p88):

- (111) the shooting of the hunters
- (112)(i) the growling of lions
- (ii) the raising of flowers

Chomsky 认为(111)的歧义必须在转换层面解释。由此需要建立下面两种转换:

$$NP-C-V \rightarrow the-V+ing-of+NP$$
$$NP1-C-V-NP2 \rightarrow the-V+ing+of+NP2$$

$NP-C-V \rightarrow the-V+ing-of+NP$	$NP1-C-V-NP2 \rightarrow the-V+ing+of+NP2$
lions growl $\rightarrow$ the growling of lions	John raises flowers $\rightarrow$ the raising of flowers
the hunters shoot $\rightarrow$ the shooting of the hunters	they shoot the hunters $\rightarrow$ the shooting of the hunters

于是,the-V+ing+of+NP 这样的序列不是由短语结构生成的核心句。

Chomsky(1957)用转换解释歧义是以不启用语义格或题元角色为前提的,所以必须使用转换才能解释清楚这里的歧义,即两个不同的语素序列可以转换成相同的语素序列。实际上,如果我们启用语义格或题元角色,说上面的 hunters 可以是施事,也可以是受事,这时就可以不用转换,因为 the shooting of the hunters 语素序列都是更短的语素序列的直接线性组合,是可以

划分层次的：

(the) ([shooting] [of the hunters])

即通过线性叠加就能生成这样的组合：

shoot + ing → shooting

hunter + s → hunters

the + hunters → the hunters

of + [the hunters] → of the hunters

[shooting] + [of the hunters] → shooting of the hunters

(the) + (shooting of the hunters) → the shooting of the hunters

至于语义上的歧义，可以用语义格来说明，这是另一种解决问题的办法，可以不用转换规则。可见这里引入转换也不是转换必要性的理由，而是在不考虑题元角色的情况下，用转换可以更好的解释歧义。所以我们说 Chomsky 在这个实例中对转换的论证也是一个描写的优化问题。

Chomsky(1957, 8.2, p88)用另一个实例论证了转换对解释歧义的作用。我们仍然可以通过线性程序来处理这个实例。Chomsky 的实例是：

I found the boy studying in the library

这个句子有两个含义，一个是说发现这个孩子在图书馆学习，一个是说发现了在图书馆学习的孩子。Chomsky 认为这个句子是由两个核心句转换而来的：

I found the boy

the boy is studying in the library

由于转换的方式不同,形成了歧义。

我们仍然可以从线性分析程序来生成这个句子,因为这个句子是可以分层次的:

(I)([found][the boy studying in the library])

至于歧义的形成,是因为语类不同:

(_{NP}I)(_{VP}[_Vfound][_{NP}the boy studying in the library])

(_{NP}I)(_{VP}[_Vfound][_{VPing}the boy studying in the library])

这里的实例翻译成汉语有比较大的区别:

NP—Verb—NP (我发现了在图书馆学习的孩子)

NP—Verb—S (我发现这个孩子在图书馆学习)

因此,用这个实例论证转换层面的存在理由,仍然是一个优化问题。再考虑 Chomsky(1957, 8. 2, p89—90)的例子:

John was frightened by the new methods.

句子有两种意思,一种意思是新方法吓住了 John,一种是用新方法吓唬 John。这两种意思可以分别体现在下面的句子中:

the picture was painted by a new technique

the picture was painted by a real artist

Chomsky 认为不同的含义都是转换造成的。其实从线性程序看,这些句子也都是可以进行层次分析的,不同的意义主要在于 by 后面引入的题元角色不同。

5.4 转换规则的必要性

下面我们要证明的问题是,转换并不仅仅是在某些方面有优越性的问题,而是生成程序的必要层面。我们的讨论除了围绕生成语法展开,也涉及结构语言学理论中和转换有关的问题。

由替换、扩充、叠加构成线性组合程序,这一程序生成的组合都可以进行层次分析。反过来说,凡是能够进行层次分析的组合,都是线性程序生成的。只要我们能够证明有的组合不能进行层次分析,就能说明有的组合不是线性程序生成的,必须要引入转换。

为了便于讨论,先需要清理一下转换的实质。传统语法中也有关于固定格式的说法,比如汉语:

连 X 也 Y (连马也吃、连残疾人也欺负)

越 X 越 Y (越看越清楚、越长越大)

其他语言中也有许多类似的用法。Goldberg(1995,2006)开始从理论上解释固定用法,提出构式(construction)的概念,认为构式是不能直接从组成部分预测的模式。在上面给出的汉语实例中,X 和 Y 都是可以进行平行周遍替换的,因此构式仍然属于规则组合,和一般的结构关系不同的是有独特的语法意义或构式意义。

还有一些特殊格式有一定的出现条件。比如:

一锅饭吃两个人	两个人吃一锅饭
一张床睡三个人	三个人睡一张床

这里表示一种分配关系,语法意义是也可以说是由构式决定的,动词进入这样的构式有一定的条件。又比如:

来了一个人
跑了一只猫
下了一场雨

这也可以看成是一种构式,语法意义表示出现和消失。

很多不能用替换生成的格式,也可以看成是一种构式,比如汉语的把字句、被字句等。因此处理生成句子的思路还有另一种方式,即在承认替换的同时,承认构式,不讨论构式是怎么在短语规则的基础上形成的,即不讨论转换的生成程序,只列出构式。但是即使这样,也还是要解释构式的语义结构的关系。比较:

吃饭
把饭吃了
把饭吃了,碗洗了。

仍然需要解释“饭”在这里都是受事,同时要承认后两个实例不是通过替换规则形成的。第2个句子是“把”字虚化后形成的,是一个历时过程,第3个句子是删除“碗”前面的“把”形成的,是一个共时过程。可见,不能用替换生成的组合是由历时和共时两方面造成的。从某种意义上说,构式理论是只承认既成的格式,只解释这些格式的特殊语法语义,不研究这些格式的形成机制。如果要追问这些构式是如何形成的,还是会涉及转换问题。

下面我们在论证转换的必要性时,有一个很重要的出发点,

即转换的引入是针对前面讨论的线性分析程序在生成句子时不具有充分性,即有些组合不可能通过线性程序生成,必须要用转换才能生成。如果先暗中用移位、添加、删除、变化等手法修正或扩展替换的概念,或者说把转换式用特定的构式接受下来,然后说自然语言的描写可以不用转换,这实际上是暗中接受了转换,因为替换、叠加等线性程序是不包含移位、删除等程序的。

5.4.1 英语助动词和转换

Chomsky(1957,p38)同时用短语结构规则和转换规则对下面的句子进行了分析:

the man has been reading the book

Chomsky(1957,5.3,p39)给出了描写 auxiliary verb 的规则:

- (28)(i) $Verb \rightarrow Aux + V$
 (ii) $V \rightarrow hit, take, walk, read, etc.$
 (iii) $Aux \rightarrow C(M)(have + en)(be + ing)(be + en)$
 (iv) $M \rightarrow will, can, may, shall, must$

(29)(i)

$$C \rightarrow \left\{ \begin{array}{l} S \text{ in the context } NP_{sing}^- \\ \emptyset \text{ in the context } NP_{pl}^- \\ past \end{array} \right\} i$$

- (ii) Let Af stand for any of the affixes $past, S, \emptyset, en, ing$. Let v stand for any M or V , or $have$ or be

(i. e. , for any non-affix in the phrase *Verb*).
Then:

$$Af + v \rightarrow v + Af \# ,$$

where # is interpreted as word boundary.

- (iii) Replace + by # except in the context $v - Af$. Insert # initially and finally.

(28)属于短语结构规则,箭头左边只有一个符号,(29)属于转换规则,有的箭头左边有两个符号。Chomsky 认为如果完全用短语结构规则处理助动词和动词的形态变化会相当复杂(It appears to be quite complex if we attempt to incorporate these phrases directly into a $[\Sigma, F]$ grammar)。

我们认为这不是一个复杂的问题,而是必要性问题,即短语结构语法不可能独立生成上面(29)所要生成的形式。因为从叠加和扩展替换都不能生成这样的句子,即下面的层次分析不能进行到底:

[the man]([has been reading][the book])

其中的[has been reading]不能再往下作层次分析。再考虑下面的层次分析,再往下分也会受阻:

[the man] / [has[been [reading the book]]]

5.4.2 主动和被动的关系

Chomsky(1957,p42)认为被动句的描写也是一个简单性和复杂性问题。如果在短语结构语法中描写主动句和被动句,比如:

John admires sincerity

Sincerety is admired by John

则每个主动句和相应的被动句都需要短语结构规则生成两种结构:

NP1—V—NP2 (主动句)

NP2—*is*+Ven—*by*+NP1 (被动句)

Chomsky(1957, p43)认为用下面的转换规则更简单:

(34) If S1 is a grammatical sentence of the form

NP1—Aux—V—NP2

then the corresponding string of the form

NP2—Aux + *be* + *en*—V—*by* + NP1

is also a grammatical sentence

Chomsky(1957)没有进一步解释为什么转换在这里更简单。我们进一步来讨论这个问题。这需要从生成规则上来理解。如果只用短语结构规则,这两种结构需要下面的规则(有的步骤省略):

$S1 \rightarrow NP1 VP1$

$S2 \rightarrow NP2 VP2$

$VP1 \rightarrow Aux - V - NP2$

$VP2 \rightarrow Aux + be + en - V - by + NP1$

如果用包括短语结构规则和转换规则的语法,需要下面的规则:

(1) $S \rightarrow NP1 VP$

(2) $VP \rightarrow V NP2$

(3) NP1 - Aux - V - NP2 $\rightarrow$ NP2 - Aux + be + en - V - by + NP1

包括转换规则的语法显然少一条规则,而且已经显示了内部 NP1 和 NP2 的相互关系。更进一步说,VP1 还有很多时态对应形式,如果 VP2 要对应起来,还会增加规则。由此可见转换规则的优化。

在汉语中情况也相似,主动句和带“被”字的被动句可以描写为:

(1) S $\rightarrow$ NP1 VP

(2) VP $\rightarrow$ V NP2

(3) NP1 V NP2 $\rightarrow$ NP2 被 NP1 V^①

这是带有转换规则的描写。这三条规则可以生成语义相关的主动句和被动句。(1)和(2)是短语结构规则,可以生成主动句式。(3)是在(1)和(2)的基础上引入的转换规则,可以生成被动句式。

再比较不用转换规则而只用纯短语结构规则对语义相关的主动句和被动句的描写:

S1 $\rightarrow$ NP1 VP1

S2 $\rightarrow$ NP2 VP2

VP1 $\rightarrow$ V NP2

VP2 $\rightarrow$ 被 NP1 V

① 为便于阅读,在变换式中我们不用数字标码,直接将符号排列起来。另外,V 前后都得有一些修饰成分。由于我们只比较句型的描写,所以不考虑终端符号的改写规则。

由于这里只依靠短语结构规则, VP 必须有两个, 即 VP1 和 VP2, 而且 VP1 和 VP2 必须分别用两条短语结构规则描写。因此整个描写必须用 4 条规则。和前面带有转换语法的规则相比, 多了一条规则, 这是纯粹用短语结构规则描写主动句和被动句的第一个弱点。当然, 一种好的语法不能仅仅根据局部描写使用了较少的规则为依据, 还得看整个语法是否简单经济。转换规则从自动化处理的角度看比短语结构规则复杂。

我们认为这里的关键在于被动句在英语中不仅仅是一个简单性问题, 因为考虑到动词的各种变化, 被动句不可能通过线性程序生成。比较下面的实例:

The work is done by John
The work has been done by John
The work is being done by John
The work has being done by John
The work would have been done by John
The work will be done by John
The work was done by John

还可以列举出更多, 其中有的是不能进行层次切分的。比如:

(The work)(has been done by John)

其中 been done by John 并不是合法片断。也就是说, done 前面的 Aux 和 be 是作为一个整体在变化, 并不是在 done 前面层层叠加得到的。

5.4.3 线性组合程序和不连续直接成分

Wells(1947)讨论过不连续直接成分。考虑下面实例:

不连续直接成分	连续直接成分
a [better] movie [than I expected]	this movie is [better than I expected]
an [easy book] to [read]	this book is [easy to read]
[too heavy] a box [to lift]	this box is [too heavy to lift]
[His father], according to John ,[is the richest man in Scarsdale]	According to John, [his father is the richest man in Scarsdale]

Wells(1947,p55)不连续直接成分的定义是：

A discontinuous sequence is a constituent if in some environment the corresponding continuous sequence occurs as a constituent in a construction semantically harmonious with the constructions in which the given discontinuous sequence occurs.

Chomsky(1957,5.3,p41)提到短语结构语法不可能处理不连续直接成分,但没有论证。我们前面在讨论连接、替换、扩充、叠加时已经提到,由这些线性程序生成的序列都是有层次的,都可以切分出直接成分,并且这些直接成分只能是连续的。这些程序不可能生成出不连续直接成分,因此,如果存在不连续直接成分的序列,就证明自然语言中有些序列不是线性程序生成的。

Wells 的不连续直接成分实际上都有一个对应的连续直接

成分。更具体地说,从上面的实例看,Wells 的不连续直接成分实际上都对应一个可以观察到的连续直接成分,不连续直接成分的形成是因为相对应的连续直接成分的某个部分被挪到了另外的位置,Wells 没有提到转换,但移位或挪动的本质就是转换。不连续直接成分的存在证明生成自然语言的句子除了需要有线性程序,还需要有转换程序。

汉语中也存在不连续直接成分。考虑下面汉语的例子:

[毕]了[业] 毕业了

“毕了业”也不能通过线性程序生成出来,可以看成是在连续直接成分“毕业了”的基础上通过转换形成的不连续直接成分。再考虑汉语的实例:

拿出来一个苹果
拿一个苹果出来
拿出一个苹果来

第一个句子可以分析成连续直接成分,第二个句子有两种可能的分析方法,第三个句子只能分析成不连续直接成分,因为下面两种分析方法都存在问题。

[拿出][一个苹果来]
[拿出一个苹果][来]

在第一种分析方法中,“一个苹果来”是不成立的组合,在第二种层次分析中,两部分成了连谓关系,和原意不同,原意“出”和“来”是直接成分关系。所以这个实例应该分析成:

拿出来[一个苹果]→拿出[一个苹果]_i来[t_i]

即“一个苹果”移动到“出来”之间,使“出来”成了不连续直接成分,并在“来”后面留下了语迹 t_i。显然,不连续直接成分是不能用线性程序生成的。承认不连续直接成分就承认了转换。

Wells 在定义不连续直接成分时,规定不连续直接成分和对应的连续直接成分意义相同。实际上存在两种情况,有些经过转换后不连续直接成分和连续直接成分并不完全相同。汉语中的动趋式宾语的移动在语义上是有变化的。我们将在后面讨论这类问题。

5.4.4 并合与转换

考虑下面的句子的意思:

我知道他去过上海了

可以有三种理解,这三种不同的意思可以通过层次体现出来:

我(知道[他去过上海了])	(他去过上海了,我知道)
我([知道他去过上海]了)	(他去过上海,我知道了)
我(知道[他去过上海了]了)	(他去过上海了,我知道了)

第三种情况是两个“了”遭遇在一起,在表层融合成了一个“了”,所以第三种情况的层次必须还原。融合的生成过程也需要看成是转换的结果,因为这种情况不能通过替换和叠加得到,因而也不能在表层上进行层次分析的,必须要还原出深层结构或基本形式,才可以进行层次分析,才可能体现替换过程。这种情况属于规则融合,这不仅是“了”的问题,两个同音节的语素,在连读中都有可能融合。再比如:

他正在看电视

他在家

两个句子合在一起：

他(正在[在家看电视])→他正在家看电视

两个相邻的“在”融合成了一个“在”。这一类规则情况的融合，不必作为并合语素存放在语符库中，否则语符库的容量会剧增。这类规则融合可以在转换规则库建立一条规则：

$A + A \rightarrow A$

5.4.5 非线性组合

不连续直接成分的前提是说存在一个可观察到的连续直接成分，由于某个成分的移动造成不连续直接成分。在自然语言中，还有一些序列，用扩充和叠加不可能生成，比如汉语：

将他的军

下面的层次切分都不成立：

* [将][他的军]

* [将他的][军]

和不连续直接成分不同的是，这些非线性组合并没有一个对应的连续直接成分。汉语母语者知道将军的对象是“他”，但下面线性结构都不成立或相当勉强：

? [将军][他]

? [对他][将军]

? [向他][将军]

因此,这类非线性组合并不是从某个可观察的线性组合转换过来的,而应该是从某个更深层次的抽象组合转换过来的。这是我们后面要讨论的必须承认深层结构的理由之一。下面的非线性组合也不是从可观察的线性结构经过移动形成的:

研不研究/道一个歉/跳得过/写封信

汉语中这类非线性组合并不多,英语中比较多。比较英语句子:

Is he speaking English?

Does he speak English?

Has he spoken English?

Did John break the cups?

What did John break?

Who broke the cups?

The cups were broken by John.

Were the cups broken by John?

What was broken by John?

从层次分析的角度看,要切分出 Does she speak English 的直接成分,会切分出 she speak English 这样不合语法的片断,因为 she 是单数第 3 人称,而 speak 不是。在该实例中,甚至考虑不连续直接成分也不能解决问题。从句子生成的角度看,这样的

句子不是扩充或叠加生成的,因为减缩替换(扩充的逆过程)和递减(叠加的逆过程)都不能还原出基式。英语还不是典型的屈折语,已经有大量的组合不能由线性程序生成,其他印欧语在这方面可能遇到的问题会更多。

英语中很多基本句子,用线性程序都不能生成。比如,述宾结构中的层次划分是相当勉强的:

speak English
 speaked English
 has speaked English
 has been speaking English
 is speaking English

这类组合的生成过程用扩充来解释很困难,用叠加来解释也很困难。这可以说明英语中使用转换生成句子的情况比汉语更多。

Chomsky 提出短语结构规则,基本结构是 $X \rightarrow YZ$,即一个结构由两个直接成分组成。前面提到,Chomsky 的纯形式规则中还有下面的短语结构规则(Chomsky,1957,p30):

$Z \rightarrow ab$
 $Z \rightarrow aZb$

这一规则可以构成:

ab
 a[ab]b
 a(a[ab]b)b
 a{a(a[ab]b)b}b

注意在这里引入了递归的重写规则：

$$Z \rightarrow aZb$$

递归形式的本质是箭头左边的非终端语符可以在箭头右边出现。这个规则的本质是插入。这样的例子在汉语中可以观察到：

找裁判员麻烦

找足球裁判员麻烦

实际上这类情况不可能通过对基式的替换或叠加生成，其生成机制已经超出了直接成分理论。所以，Chomsky 所提到的短语结构语法实际上是比较直接成分理论要强的理论，已经包含了部分非连续直接成分的因素。当然，如果承认直接成分中有不连续直接成分，或者承认插入替换，直接成分理论和这里的短语结构模型就是等价的。问题是，插入本身已经不是线性程序。

思考问题

1. 广义短语结构语法能否从根本上取消转换？
2. 前人在论证转换的必要性上存在什么问题？

6. 音系分析程序的非线性化

以上讨论了在句法层面转换规则的必要性,根本原因是因为连接、替换、叠加并不能生成全部的句子,有些句子必须要借助转换才能生成。在语音层面,通过语音单位的连接、替换和叠加的方式可以生成很多语音序列。

$$l+u\rightarrow lu$$

在生成音系学中可以用生成规则描写以下汉语 a 音位的规则:

$$\begin{aligned} a &\rightarrow a/_n \\ a &\rightarrow \alpha/_\eta, _u \\ a &\rightarrow A/_* \end{aligned}$$

这种生成过程本质上也可以通过叠加完成,只是要求分类更细:

$$\begin{aligned} a+n &\rightarrow an \\ a+i &\rightarrow ai \\ \alpha+\eta &\rightarrow \alpha\eta \\ \alpha+u &\rightarrow \alpha u \\ A+* &/A* \end{aligned}$$

连接前或叠加前的音段和叠加后的音段并没有变化。在汉语中,语素音形基本上都能用连接、替换和叠加程序生成。这就是说,语素内部的形式都可以通过线性方式来形成。

本章要论证,在音系层面,和词法、句法层面类似,连接、替换、叠加等线性生成手续也不能生成全部的语音序列,音系层面还需要一些派生规则,这些派生规则的性质和转换规则是相似的。所以在音系层面,也存在非线性或超线性问题。和句法层面一样,语音单位也存在规则组合和不规则组合的问题,而且把哪些形式看成线性组合,哪些看成变音规则,需要考察平行周遍条件。

主要问题集中在对基本单位和规则的认识。语素基本音形(语素基形)是最重要的研究起点。首先必须认识到语素内部音形单位及其组合规则是一种方式,语素基形之间的组合规则和音段的分布是另一种方式。这里所说的语素基本音形和通常所说的语素本音不同,语素本音是指历史上比较早的形式,而语素基本形式是指不可推导的形式。语素基本形式的严格定义后面将给出。确定语素基本音形必须考虑单位和组合的关系。结构语言学没有认识到语素基本音形的重要性,生成音系学把语素和词都作为基础,开始重视语素在音系分析中的重要性,但语素音形和词音形的关系讨论不充分。规则问题生成音系学认识到了,并充分讨论了规则的生成机制,但没有提出判定规则和不规则的标准。

生成音系学区分可预测(predictable)和不可预测来确定深层语音形式(Chomsky and Halle, 1968, p12)。词库中存放的应该是项目(items)的个别特征(idiosyncratic properties),即不可被普遍规则预测的特征。比如 cat 的复数 cats 是可预测的, I 的复数 we 是不可预测的。这对哪些形式需要存放在词库中是有价值的,但可预测的说法还不完全反映语音规则的本质,如何判定可预测也需要研究。汉语普通话“一”从 yi^{55} 出发, yi^{51} 是可以预测的。从 yi^{51} 出发, yi^{55} 也是可以预测的。上声从 214 出发可以预测 35、21 的分布,从 21 出发也可以预测 214、35 的分布。汉语普通话“一”的变调和上声变调都是可预测的,但平行周遍

情况不一样,后面我们将进一步展开分析。我们认为问题的关键在于怎样确定变音规律或音变规则。变音规律和音变规律相似,只不过音变规律是历时的,变音规律主要是共时的。我们所说的变音规律,是指相同的语音基形组合必须有相同的变化。或者说,不同的变音必然有不同的基形条件。找出确定变音规则的条件才能给出判定基形的标准。所谓不同,又要涉及对立的概念。对立是相当根本的初始概念。两个语素尽管有部分相同的语音表现,但只要有一处语音表现有对立,就构成不同的语素基形。对立的语音表现可以作为基形。

生成音系学强调规则的研究,通过规则来确定深层形式,这是一个重要的进展。当然这取决于把规则用在什么层面,生成音系学在语素音系内部也使用规则,其实语素音形内部的规则与语素音形之间的规则是相当不同的。后生成音系学开始区分纯语音规则和语法语音规则,观察到了很多重要现象,但跟区分语素音形内部的规则和语素音形之间的规则仍然不同。

到此又把我们引回到了结构语言学的发现程序,单位和规则的关系再一次提到日程上来。语言发现程序或分析程序是田野调查不能绕开的问题,音系中核心语音单位或基本语音单位的提取程序也是不能绕开的问题。关键是发现程序从哪个层面开始?核心语音单位是只能区别语素音形并且不能通过规则推导出来的语音单位,是音系分析的基础。但是后结构语言学发现程序从比语素更长的词或言语片断直接提取音位,再从音系到语素再到词和话语展开描写,会遇到困难。要提取核心语音单位必须先提取到语素音形。语素音形的数目对不同的调查者来说是唯一的,音位数目和音位规则对不同的调查者来说不是唯一的,所以语素音形是音系分析的起点。

要弄清问题的实质,让我们先从整体上分析一下当时的理论背景。结构语言学发展到后期试图提出一套发现程序,先描写语音,再描写语素,最后描写语法。Hockett(1942)提出音系

描写不能参考语法层面的信息,以免有循环论证。所谓不参考语法层面的信息是指不能先假定语素、词等单位已经获得。所谓循环论证是指提取音位时依赖语素、词等概念,提取语素、词时又依赖音位等概念。Harris 的 *From phoneme to morpheme* (1955) 和 *From morpheme to utterance* (1946) 两篇论文体现了这种从音系到语素再从语素到话语的思路。生成音系学反对研究发现程序,认为音系描写需要参考语法语义因素(Chomsky, N. M. Halle and F. Lukoff, 1956)。生成音系学参考语法语义因素的说法有时候会引起误解,其本质含义应该理解成音系基本单位的确定应该考虑词界、构词环境等条件,这一点在 Halle (1959) 的论述中是比较清楚的。Halle (1959) 还认为系统音位 (systematic phoneme, 相当于语素音位 morphophoneme) 才是音系的基本单位,而结构语言学的音位 (phoneme) 是多余的。Halle 还认为在确定系统音位时,一个语素只需要确定一个基本形式,其他变化形式通过派生来确定。Halle 已经注意到了音系基本单位的一些本质特点,留下的两个问题是:

1. 按照 Halle 的思路,如果没有词界的信息,下一步确定基本单位就无法进行。但词界在很多语言中并不是很清楚的,比如汉语。在词界不清楚的情况下,怎么确定音系基本单位? 英语中所有格 s 前的成分是可以扩展的,也是一个问题。比如 the king of England's (英国国王的)。另外,很多规则并不受词界限制,比如汉语上声变调,“我有”和“雨伞”变调模式是相同的。

2. 一个语素只需要一个基本形式的说法还不明确,并没有严格区分周遍规则和不周遍规则。比如汉语的上声语素,只确定一个基本形式没有问题,但“一”只确定一个基本形式就很困难。“亚”的情况又是一种。

问题的核心是音系分析从什么地方开始。Chomsky and Halle (1968, p7)是在 formative 语形(相当于语素)序列(lexical representation)的基础上派生出词音形序列(phonological representation),再派生出音值。这是一种很重要的思路,词库中储存的都是 formative 的音形,这就保证了纯区别层面的音形。但 Chomsky and Halle 没有论证为什么要以语素音形为基础。在 Chomsky and Halle 的理论中,词音形也是一个重要的层面。语素音形和词音系到底是什么关系?实际上以上存在的两个问题仍然没有得到解决,语素音形存在的理由没有得到充分论证,语素音形的重要性一直没有得到足够的重视。问题又回到了发现程序的论证上。至今很多人以为,只要在田野调查中记录足够的故事和词、词组,回到研究室就可以提取音系。我们的问题是,我们在田野调查中总要面临调查步骤的问题,总要有一个有效的办法来完成语言调查,这本质上就是发现程序的问题。儿童在很短的时间能够迅速学会自己的母语,也有一个学习过程,研究这样的过程实际上也是在研究发现程序。所以发现程序的研究是语言研究非常重要的工作。田野调查程序是观察和检验语言学习机制理论的重要窗口,如果一种语言理论不能说明田野调查程序,很难说这种理论已经解释了语言学习机制。生成语法由于没有回答发现程序的问题,所以在田野调查中生成语法的理论会遇到一些困难。结构语言学对发现程序的研究反映了结构语言学后期对研究目标提出的更高要求,只是从语音到语素再到语法这个程序本身存在问题。

根本问题就在于对语素和词的看法。后结构语言学要找发现程序,而词是不容易判定的,所以不在音系分析中起用词的概念,连语素音形的分析也是在音位分析后才出现,因为确定语素还得参考意义。生成语法预先假定词和语素,但不能说明发现程序,所以通常都是以俄语、英语等词界比较明确的语言出发进行分析比较成功。汉语中,词的判定比较困难(徐通锵,1997;陈保亚,1999)

其实早期结构语言学家 Bloomfield(1933)在分析音位的时候所列举的例子都是单音节词,实际上也是单音节语素。Bloomfield 当时没有论证为什么要这样做。A. Martinet(1960)主张由单音节词入手分析音位。王理嘉(1988)认为音位分析最好从通常称之为“词”的成分入手,而且最好是从简单词干形成的单音节词入手,并用汉语普通话的例子加以论证,认为普通话的音位分析应当从单音节语素入手。王洪君(1999)主张从汉字音的角度出发分析汉语音位,更容易解释音系的系统性和演变特点。我们认为 Bloomfield(1933)的具体做法是有一定道理的,王理嘉(1988)和王洪君(1999)的观点也是有道理的,不过都还没从田野调查的发现程序入手考虑问题,所以得到的结论主要认为从单音节词、单音节语素或单字音出发进行描写对解释音系的系统性和演变机制比较方便有利,所以描写的起点问题在当时还是一个好坏问题。至于描写的起点到底应该从单音节语素、单字音、单纯词、单纯词干入手还是语素入手,它们之间有没有区别,还需要从分析程序的角度进行更多的论证,因为我们在长期的田野调查中注意到,选择什么单位作音系描写的起点,是音系分析能否充分进一步展开的可行性问题,需要我们有一个明确的选择和论证。考虑到汉语和中国境内民族语言以及其他相关语言在这方面存在的共同问题,尤其是田野调查中所必须面临的分析程序问题,我们将给出统一的论述,不仅仅考虑对解释音系结构和音系演变有利与否的问题,而且要尽可能找出音系描写所必须依赖的程序。我们的基本结论是:音系分析的基点应该建立在语素音形上(不只是单音节语素)。整个论证重点集中在这样一条思路:音系的描写首先要提取核心语音单位,而提取核心语音单位必须先假定语素基本音形已经被提取,而提取语素基本音形必须和提取语音规则同时进行,从词、话语片断、文本入手描写音系都会遇到困难。

6.1 自主音系学的不充分性

结构语言学试图在不考虑词法和句法层面信息的前提下，对音系进行独立描写，即直接从话语片断的差异中提取出对立的音位。这一分析方法被称为自主音系学，即不考虑语法层面信息，直接从纯语音层面对音系进行自主分析。生成音系学分析了自主音系学的缺点，认为这种分析方法会使音系结构相当复杂。

我们认为自主音系学所遇到的还不仅仅是一个复杂问题，不仅仅是一个优选问题。如果严格地从纯语音层面分析音系，自主音系学的分析程序是不充分的，即不能对音系进行充分描写。我们可以从音系对语素的依赖上来认识这一点。如果没有语素的前提，很多语音规则是找不到的。比如，从纯语音记录出发我们可以记录到汉语普通话单音节的 4 个单音节调型 55、35、214、51 和 15 个双音节调型（还没有包括轻声）。双音调型的分布是有空格的：

55+55	55+35	55+214	55+51
35+55	35+35	35+214	35+51
214+55	214+35		214+51
51+55	51+35	51+214	51+51

仅仅靠语音层面的信息，我们只能断定 214+214 在表中不出现。这就存在两种情况：

- 1. 可能是不能组合。
- 2. 也可能是组合后产生变调。如果是变调，只从纯语音层面的条件也不能断定 214+214 变成了上面哪一种模式。

纯粹从语音分布出发不能解释上面的疑问。只有当上述分析和语素结合起来,具体地说,只有当我们事先提取了语素,知道了每个语素的独立调型而不仅仅是音节的调型,才能判定 214+214 变成了 35+214。比如:

语素组	前语素调型	后语素调型	组合调型
徒手	徒:35	手:214	35—214
强手	强:35	手:214	35—214
好手	好:214	手:214	35—214
老手	老:214	手:214	35—214

由此可见,在没有语素信息的前提下,不可能对汉语声调行为进行充分描写。从音位层面看,也就是不可能判定哪些音位有哪些变体。

即使双音节调型没有出现空格,也存在声调判定问题。比如四川德阳话:

55—55	55—35	55—51	55—214
31—55	35—35	35—51	35—214
51—55	51—35	51—51	51—214
214—55	214—35	214—51	214—214

其中 31—55 有一部分是 31+214 变来的,从纯语音分布上看不出这一点,从语素和语素组合的关系上就能看出来:

单字调	单字调	偏正 (名)	述宾	并列	主谓	偏正 (谓)	述补
磨 ³¹	肉 ²¹⁴	磨 ³¹ 肉 ⁵⁵	磨 ³¹ 肉 ²¹⁴				
鱼 ³¹	肉 ²¹⁴	鱼 ³¹ 肉 ⁵⁵		鱼 ³¹ 肉 ²¹⁴			
油 ³¹	漏 ²¹⁴	油 ³¹ 漏 ⁵⁵			油 ³¹ 漏 ²¹⁴		
迟 ³¹	到 ²¹⁴	迟 ³¹ 到 ⁵⁵				迟 ³¹ 到 ²¹⁴	
吃 ³¹	够 ²¹⁴	?					吃 ³¹ 够 ²¹⁴

31 调加 214 调时,如果满足下面两个条件就要变调:

1. 名词性的。
2. 偏正结构。

其他情况下不变调。这个现象说明变调会受结构关系和分布功能两种因素的影响,如果没有句法层面的信息,就不能充分描写汉语德阳话的变调行为。

6.2 语素基本音形与变音

汉语上声变调的描写说明音系分析必须先提取语素的基本形式(语素基形)。结构语言学的自立分析原则之所以不充分,就是因为没有区分语素基形和变音,而这种区别的前提是先获得语素基形。我们所说的语素基形是指不包括语素规则变音的语素音形。比如“手”的基形是 shou^{214} , “ shou^{35} ”则是在上声前的规则变音。如果我们注意观察后结构语言学的双向唯一性原则所遇到的麻烦实例,就会发现都和变音有关系。观察下面的实例:

bit bid

bet bed

pot pod

在美国中西部英语中,右栏的元音比左栏的元音要长一些,但没有对立。Bloch(1941)的 *Phonemic Overlapping* (1941)代表了不参考语法语义的条件下独立描写音系的方法,把以上出现在两个不同条件下的 i 归纳成一个音位。根据同样理由,两个不同条件下的 e 也归纳成一个音位。但以上的 o 却归入不同的音

位。Bloch 的理由是, pod 中的 o 和 pa'd(pa would 的合音)中的 a 读音相同, 如果 pot 和 pod 两个元音归入一个音位, 就会出现同一个读音在相同的语境中(p-d), 既是 o 音位, 也是 a 音位, 违反双向唯一性原则。Bloch 这样做维持了双向唯一性原则, 但和直觉相差很大, 而且不对称。其实根本问题就是把语素音形内部的音段分析和语素组合的音段分析混在一起。

语音的对立有三个层次。在语素范围内的对立, 在词范围内的对立, 在话语范围内的对立。这三种对立是不一致的。在一种范围内对立的两个音段, 在另一种范围内可能对立, 也可能不对立。比如汉语实例:

对立音段	语素	词	话语
o:u	佛		
	弗		
-n:-r	二	二	
	摠	摠	
m:p	面	面子	有面子
	辩	辫子	有辫子
ɛ:ə		鸡儿	小鸡儿
		街儿	小街儿
a:ɑ			床边
			船边
o:ɔ		果儿	
		滚儿	

在以上材料中, 要从话语音形、词音形确定语素基本音形都是有困难的。而从语素基形出发再根据规则推导变音是可行的。后面我们会更详细的讨论这种不可逆关系。由于语素音形和各种变音都是可观察事实, 都需要记录。

有了语素基本音形和变音的区分, 可以看出双向唯一性

原则在语素基形内部是有效的。语素音形内部的音段可以通过增加变体减少音位最终满足双向唯一原则。语素音形之间的语音规则不是可逆的,或者说不满足双向唯一原则。比如,汉语普通话连上变调是语素音形之间的规则,上声如果用变调形式 35 作基本形式,就会出问题。当然有人可能说 jie 中的 e 是[ε],由于又有人把 jian 中的 a 作[ε],就会出现同一个音子被当作两个音位。不过这并不违反双向唯一性原则,因为这是同一个音子出现在不同的语音条件中,并不是出现在相同的语音条件中,这和 35 出现在相同的 214 中而被归为两个调位并不一样。

事实上,我们并不会在语素基本音形分析中把相同语音环境下相同的音作为不同音位,这会违反同一性原则,使分析无法进行。相同语音环境下不同的音如果不区别意义都可以看成相同的音,更何况相同的音。我们认为,在语素基本音形内部,坚持双向唯一性原则就是坚持同一性原则。

6.3 音系分析条件:语素基本音形

结构语言学试图在完全独立于语素层面的情况下通过分布来提取音位,这是发现程序第一步。Chomsky(1964, 4.3)把这种方法称为分类音位学(Taxonomic phonemics)。Chomsky(1964, p80)归纳出分类音位学的音系表达的四个条件:

1. linearity
2. invariance
3. biuniqueness
4. local determinacy

其中“双向唯一性”(biuniqueness condition)条件是结构语言学

音系学理论的一条重要原则。这一术语是 Chomsky (1964, p80) 正式提出的。更早的时候, Halle (1959) 已经对结构语言学的这一原则进行过批评。Chomsky (1964, p80) 曾对“双向唯一性”原则有个一个比较严格的表述:

严格地说, 双向唯一性条件断定每个语音序列都由唯一的音位序列表达, 并且每个音位序列都表达唯一的语音序列。(Technically, the biuniqueness condition asserts that each sequence of phones is represented by a unique sequence of phonemes, and that each sequence of phonemes represents a unique sequence of phones.)

这条原则实际上要求在特定的语音条件下, 一个音位对应于一个音子, 一个音子对应于一个音位。为了更进一步分析双向唯一性原则, 我们再概括一下双向唯一性原则的两个要点, 便于后面讨论:

1. 同一音位的不同音值必有不同的语音条件(相当于互补原则)
2. 同一音值归并成不同音位必有不同的语音条件

对于第 1 条, 比如汉语 /a/ 音位, 在 pan 中对应于 [a], 在 pang 中对应于 [ɑ], 在 pa 中对应于 [ʌ]。对于第 2 条, 考虑:

jie⁵⁵ 街

juan⁵⁵ 娟

第 2 条肯定会涉及审定音值的问题。如果认为 e 和 a 读音相同, 可以说相同的音值在不同的条件下归属不同的音位。如果

认为 e 和 a 读音不同,符合不同的音值可以归入不同的音位。

但下面实例违背了双向唯一性原则的第 2 条,同一音值在相同的条件下归并成了不同的音位,比如成都话的儿化词中:

干儿 $\text{kanr}^{55} [\text{kər}^{55}]$

根儿 $\text{kən}r^{55} [\text{kər}^{55}]$

从 Chomsky(1964, p80)对双向唯一性原则的定义看,以上实例并不满足每个音子序列用唯一的音位序列表达(each sequence of phones is represented by a unique sequence of phonemes), $[\text{kər}^{55}]$ 实际上是用两个音位序列表达,即 $[\text{kanr}^{55}]$ 和 $[\text{kən}r^{55}]$ 。

Chomsky(1964, p80, note 15)还有一个关于双向唯一性原则的弱定义:

每个音位序列代表了一个唯一适合于自由变体的音子序列(Each sequence of phonemes represents a sequence of phones that is unique up to free variation.)

前面实例是两个音位序列表达了同一个音子序列,所以 $[\text{kər}^{55}]$ 不是唯一的。

我们可以从 $[\text{kanr}^{55}]$ 这一音位序列推导出 $[\text{kər}^{55}]$,从 $\text{kən}r^{55}$ 推导出 $[\text{kər}^{55}]$,但反过来,在不考虑语素意义的差异的情况下,却不能从 $[\text{kər}^{55}]$ 推导出确定的音位序列。这说明,如果分析语音只从语音层面入手,即通过纯语音的发现程序,不可能发现“干儿”和“根儿”在内部构造上的差异。我们前面提到的普通话上声变调,也存在这样的问题。生成音系学批评双向唯一性原则时认为,从纯语音条件分析中得到音位的发现程序是行不通的,音系学的研究思路应该相反,走生成的思路,即通过深层单位和规则生成具体的音值。这里的关键是

从底层到表层的推导要求。如果不从成都话的[kər⁵⁵]去推导什么,直接说[kər⁵⁵]有“干儿”和“根儿”两种意思,也是一种描写方法。这样一来可以保住双向唯一性。从这个角度看,生成音系学对结构语言学的批评是不充分的,这一批评只说明了纯语音层面的发现程序不能有效说明底层单位的还原和规则的关系,并不证明发现程序是完全行不通的。生成音系学通过深层单位和规则生成语音表达,这只是语言学习完成后总结出的结果,只是语言调查结束后总结的结果,在语言学习和语言调查的开始阶段,接触的都是具体的材料,然后通过材料的分析得到单位和规则。深层单位是怎么来的,仍然有一个发现程序的问题。

问题不是出在发现程序本身,而是出在结构语言学音位的发现程序完全被限制在纯语音层面,语素基本音形和语素组合后形成的音形没有得到区分。结构语言学的音位发现程序是不首先提取语素基本音形的发现程序。事实上,在语素组合过程中,不同的音位可能变成相同的音,比如成都话的“干儿”和“根儿”,出现所谓相同的音子对应不同音位的现象,即[kər⁵⁵]中的/ə/分别对应于[kanr⁵⁵]中的/a/和[kən⁵⁵r⁵⁵]中的/ə/。汉语普通话上声变调也属于这样的情况,比如/t'u³⁵ kai²¹⁴/中的35分别对应于/t'u³⁵ kai²¹⁴/(涂改)中的35和/t'u²¹⁴ kai²¹⁴/(土改)中的214。Bloch(1941)分析 pod、pa'd、pa也属于这种情况,pa'd的读音和 pod 的读音相同,是语素组合的结果。如果我们能够先切分语素,提取语素基本音形,再从语素基本音形中获得基本语音单位,就可以得到深层单位和规则,从基本音形和规则生成可观察到的表层语音。在语素基本音形内部,双向唯一性原则是可以实现的。当然还必须解决自由变体问题。

更具体地说,双向唯一性原则之所以遇到困难,从根本上说就是因为没有区分语素音形内部的语音行为和语素组合音形的

语音行为。为了更进一步认识双向唯一性原则的本质,我们再来分析这一原则的最初含义。双向唯一性原则最早是在 Bloch (1941)讨论音位的交叉分布(intersection 或 overlapping)涉及的。Bloch 当时区分了两种情况,部分交叉分布(partial intersection)和完全交叉分布(complete intersection):

如果某个语音 *x* 在某一组语音条件下被归为音位 A,而同样的语音 *x* 在一组不同的语音条件下被归为语音 B,音位的这种分布可以叫作部分交叉;如果 *x* 在相同的条件下出现有时候被归为音位 A,有时候被归为音位 B,音位的这种分布就是完全交叉。

对于部分交叉,汉语中 *ian* 和 *ie* 的韵腹就是实例, /a/ 和 /e/ 是两个不同的音位,但在 *ian* 和 *ie* 两个不同的语音环境中都读 /ɛ/, 即同一个 /ɛ/ 在 *ian* 中被归在 /a/ 音位, *ie* 中被归在 /e/ 音位,形成部分交叉分布^①。Bloch 认为部分交叉不会影响音位分析原则,在实际中也不会带来不确定性。至于完全交叉分布,我们认为 Bloch 提出这一概念本质上是针对中和来的(neutralization)。我们先来分析 Bloch 给出的一组实例。为了表述更直观,按照通常的音位分析方法,我们可以把 Bloch 这组实例分成 A、B 两组:

A	where at	at home			
B	cat	atone	about	come	rush

根据这种分组, A 组中斜体字母是同一音位, B 组中斜体字母是同一音位。A 组音位分析把 *at home* 中的 *a* 和 *where at* 中的 *a* 当做一个音位,因为这两个环境中的 *at* 实际上是同一个

① 当然,如果认为 *ian* 中的 *a* 和 *ie* 中的 *e* 不同音,就不存在 *a* 和 *e* 的交叉分布问题。可见交叉分布和判定音值的同一性是有关系的。

词。Bloch 认为这样的音位分析会违背双向唯一性原则,因为 A 组 at home 中的 a 和 B 组 atone 中的 a 读音是相同的,实际上是出现在-t 这一相同语音条件下的同一个读音,如果把这两个环境中的 a 分别归在 A 组音位和 B 组音位,则相同的读音在相同的条件下被归入了不同的音位,违背了我们上总结的双向唯一性原则的第 2 条:同一音值归并成不同音位必有不同的语音条件。

为了避免完全交叉分布,即不违反双向唯一性原则,Bloch 主张把 at home 中的 a 的音位和 atone 中的 a 音位看成同一个音位,这实际上是把 at home 中的 a 不归入 A 组而归入 B 组,这就必然形成这样的分组:

A	where at					
B	cut	atone	about	come	rush	at home

这种处理办法从根本上看就是不考虑语素层面的信息。如果从语素层面看,atone 是一个语素或一个词,at home 是两个语素或两个词,at home 和 atone 中的 a 同音,是变音的结果。上面第一种分组的前提是,我们能够判定 where at 和 at home 中的 at 是同一个词。Bloch 的分析方法是不考虑语素或词的同一次性,只从纯音值来考虑同一性。Halle(1959)、Chomsky(1964)已经从优选的角度指出自立音系学在描写音系上的复杂性,我们前面进一步论证了自立音系学的不充分性,证据是不能获得普通话上声变调行为,因此,音系分析必须以语素音形为条件。至于交叉分布问题,我们可以允许不同的音在变音过程中出现同音,这并不是语素基本音形的同音。

Bloch(1941)的分析是要走一条音系分析独立于词和语法的分析道路,Harris 后来进一步把这种思想系统化。Halle(1959)对结构语言学自立音位学不依赖语法的方法进行了批评,Ferguson(1962)为结构语言学的方法辩护,并正式提出了音

系自立(the autonomy of phonology)的概念,认为音系是完全独立于句法和形态的,并认为双向唯一性原则和局部确定(local determinacy)条件是合理的。

Chomsky(1964)对结构语言学自立音位学不依赖语法的方法再次进行了批评,提出音系是非自立的(non-autonomous, Chomsky, 1964, p106),即有些语音过程是依赖句法结构和形态结构的:

some phonetic processes depend on syntactic and morphological structure so that phonology as a whole can not be studied, without distortion, in total independence of higher level structure.

生成音系学对结构语言学的批评有一定的道理,生成音系学的音系分析确实证明要依赖词界才能说明语音行为。其实 Bloch(1941)已经对依赖词进行音系分析有所考虑。如果重读的元音无法确定,同一个词(the same word)就无法确定。比如,如果 where at 已经消失,没有办法断定 at 的重读形式, at home 中的 a 怎么办? Bloch 认为 of 就是找不到重音形式的例子。生成音系学没有回到这样的问题。我们认为,如果 at 的重读形式真的不存在,也就不存在同一个读音在相同的语音条件下分属两个音位的问题。我们面临的仍然是重读形式还存在的问题。

不过生成音系学依赖词信息进行音系分析也存在问题,在英语以外的很多语言中,包括汉语,确定词是相当困难的。从发现程序的角度看,后结构语言学家显然不允许把这样一些不清楚的概念作为研究的出发点,因为发现程序要追求严格的分析步骤。至于描写的优选问题,尽管结构语言学的发现程序可能使描写复杂化,但有实证性。这两方面可能是 Ferguson(1962)

以及其他一些结构语言学家为自立音系辩护的重要理由。尽管自立音系学有这样一些理由,但考虑到我们前面讨论过的发现程序的不充分性,自立音系学不成立。

描写的不充分性是自立音系学不成立的根本理由,音系要得到充分的描写,必须要先获得语素基本音形。后面的分析将要说明,语素基本音形是音系分析的起点,是必要条件,当然还不是充分条件,很多语音规则的确定还要考虑其他因素,如词界、语法功能、节律等,而这些都是以语素的确定为前提的。

6.4 从语素音位到语音规则和语素基本音形

音系分析必须以语素基本音形为条件,一个语素在不同的环境下可以有不同的读音,如何认识这些不同的读音的关系,是确定语素基本音形的关键。

Bloomfield(1926)区分了语音交替(phonetic alternation)、形式交替(formal alternation)。形式交替如英语复数 book-s[s], boy-s[z], ox-en, f-ee-t,这已经涉及语素音位(morphophoneme)的问题。Bloomfield 同时还认为形式交替和语音交替不同之处在于,形式交替只限于具体的形式。比如,复数交替只限于该形式,并不是所有的 s 读音都有这样的交替。Bloomfield 还认为,如果形式交替是由所伴随的音位环境决定的,就是自动交替(automatic alternation),如英语的[-s,-z,ez],如果是由其他因素决定的,就是语法交替(grammatical alternation)。Bloomfield 还区分了规则变体(regular variant)和不规则变体(irregular variant)。从 Bloomfield(1926)所讨论的问题看,已经基本上涉及了后来语素音位讨论的核心问题,只是没有用语素音位这一术语。

Trubetzkoy(1929a)首先提出了语素音位(morphophoneme)的术语。语素音位的不同形式代表了一个特定语素的变体(allomorphs)。比如俄语 luk(牧场),luga(牧场的),luk 和

lug-代表了“牧场”这一语素的两个变体。又比如英语复数语素的变体有/-s, -z, -iz, -en, -0/, 分别代表 books, dogs, horses, oxen, sheep 中的复数语素的读音。

不难看出,语素音位是语素音形的抽象形式,而语素变体是语素音形抽象形式在具体语境下的具体读音。区分语素音位和语素变体的思路跟区分音位和音位变体的思路是一致的,可以把两者的平行关系概括如下:

抽象单位	具体读音
语素音位	语素音位变体
音位	音位变体

美国结构语言学后来用语素音位来处理同一个语素具有不同表现形式这一现象。

Harris(1942)深入讨论了语素交替(morpheme alternants)和提取语素的问题,并用了 morphophonemic symbol(1. 2, 如 /najf/), morphophonemic formula(1. 2), morphophonemics(6. 1)这样的术语。Harris(1942)关于语素的定义和 Bloomfield(1933)的定义相同(Harris, 1942, Joos, 1957, p109):

每个有意义的音位序列,只要不是由更小的有意义的序列构成,就是语素。

Bloomfield(1933)对于语素的归并并没有给出明确的方法。Harris(1942, 2. 2)给出了一个确定语素及其变体的条件:

A morpheme unit is thus a group of one or more alternants which have the same meaning and complementary distribution. To make these units more similar to our present morphemes, and more serviceable for grammatical structure, we

now add a further condition: In units consisting of more than one alternant, the total distribution of all the alternants (i. e. the combined range of environments in which each of them occurs) must equal the range of environments in which some units with but a single alternant occurs.

即一个语素是一个或多个有相同意义并且互补的交替形式的集合。由于 Harris 更强调描写的形式化和严密性,所以从他那里我们比较容易看出结构语言学分析方法的某些本质。Harris(1942, 1. 1)分析了 Tübatulabal 语,可以把 Harris 的做法归纳如下:

puw	irrigate	ubuw	he irrigated
pələ • la	arrive	ə • bələ • la	he arrived
wə • ?	to pour	ə • wə • ?	he poured

Harris 认为, puw 和 -buw 是同一语素的不同形式 (different forms of the same morpheme), 这是把 puw(irrigat) 看成一个语素变体, -buw 看成另一个语素变体, ubuw(he irrigated) 看成两个语素相加, 而且每个语素都必须是可以观察到的。这种观点是一种典型的线性分布观点。

我们通过上面材料第 1 和第 2 行 pələ • la 的比较, 可以看出 Tübatulabal 语中存在一种清音浊化的变音。如果这种浊化现象在 Tübatulabal 语中是平行周遍的, 应该在规则库中设立一条规则, 而不必为很多语素设立语素变体。如果只是平行但不周遍, 处理方式可以和前面讨论的汉语“七、八”的处理方式相似, 既可以在语符库中用变体处理, 也可以在规则库中用规则处理, 并在语符库中注明哪些语素要浊化。

结构语言学家都有一种基本观点, 当 AB 两个语素组合的时候, AB 两个语素的音形一定可以在线性方向上切分出来, 而且一定要切分出来, A 和 B 都有自己独立的语音形式, A 和 B

的变体也都有自己独立的语音形式。这是结构语言学的基本思想：话语(utterance)由语素构成。

实际上话语的构成至少经过两个步骤：

- 1. 单位组合形成有层次的语素序列。
- 2. 语音规则和其他层面的规则作用于语素序列，构成话语(utterance)。

汉语多个上声字的连读变调最能证明第 2 个步骤的必要性，因为 3 个或 3 个以上的上声字组合就需要层次规则：

买[雨伞]→21-35-214
[雨伞]小→35-35-214

由于后结构语言学只承认第 1 个步骤，不承认第 2 个步骤，认为话语直接由语素构成，所以尽量通过变体分析来体现这种直接线性组合，结果使语符库很复杂，甚至不能描写话语的读音，比如 3 个或 3 个以上上声字组合的读音。如果我们承认上述两个步骤，那么在提取单位的时候，要同时考虑对比矩阵中的各种规则，以便能够有效提取单位，而不能只是简单地使用纯线性的表层对比或可观察语音对比，不顾规则的存在。

有些语素的变音有平行性但没有周遍性。Swadesh 和 Voegelin (1939)提出过一种非显露音系的理论(theory of non-patent phonology)，他们把下面名词复数的特殊变化分成两组：

A 组		B 组	
leaf	leaves	belief	beliefs
sheaf	sheaves	roof	roofs
calf	calves	proof	proofs

续表

A 组		B 组	
half	halves	hoof	hoofs
knife	knives	chief	chiefs
wife	wives	gulf	gulfs
wolf	wolves	cough	coughs
hoof	hooves	puff	puffs
life	lives		
scarf	scarves		scarfs
handkerchief	handkerchieves		handkerchiefs

Swadesh 和 Voegelin 认为存在两种清擦音语素音位 (morpho-phonemes), 可以提供两种符号来表示这种区别, 一组有确定的读音, 如 B 组的 bəlɪf, 一组的读音在复数条件下读浊音, 如 A 组中的 leaf。A 组的语素可写作 liF。Swadesh 和 Voegelin 用大写的 F 表示抽象语素音位, 可以代表 [v] 或 [f], 其具体音值根据上下文决定。比如, 如果是单数, 就是清音, 如 leaf、half、wife 等, 如果在复数标记前, 就是浊音, 如 leaves、halves、wives。Swadesh and Voegelin 说:

Different phonologic behavior requires the differentiation of the morpho-phonemes.

Swadesh and Voegelin 从根本上说还是以语素变体的方式分析问题, 根本思想仍然是语素由音位构成。更准确地说, 语素在各种条件下都有自己独立音位形式。Swadesh 和 Voegelin 的语素音位和后来 Harris 的语素音位还不完全一样, Swadesh 和 Voegelin 的语素音位实际上是语素音形中的某个音位片断, 比如 F 只是语素音形 liF 中的某个音位片断。当然, 不同时期语

素音位概念的差异不是本质的,本质的地方在于其共同性,即用变体来处理语素音形的不同读音。

我们认为,以上变音现象可以不用 Swadesh 和 Voegelin 的变体来分析,而上升到一条更普遍的规则。可以看出,上面几个语素的最后一个辅音读浊音[v]时,s 就读浊音[z],这里的浊音化其实是英语中复数化过程中存在的更一般规则的体现。比较:

mouth	mouths[mauðz]
path	paths[pɑ:ðz]
month	months[mʌnθs]
youth	youths[ju:ðz]/[ju:θs]

所以可以更进一步给出一条规则:

如果复数 s 前的辅音是浊音,s 选择[z]。

以上变音现象还可以用另一种规则来描述,即上面几个语素的最后一个辅音复数化时变成浊音的[v],可以看成是复数形式 s 读浊音引起的。因此可以有一条规则:

如果复数选择了[z],则前面直接所连接的非浊化辅音必须浊化。

这种描写实际上包含了对复数 s 的不同处理方式。过去认为复数 s 的变体是由前面的词形决定的,这种方式承认 s 本来就有清浊两个,由此决定了前面词形最后一个辅音的发音方法。这是根据对立原则。

至于哪些名词词形的收尾辅音有清浊交替,或者说/f/、

/th/后面复数在什么时候选择[z],变体描写的方式和规则描写的方式都不能确定,因为两种选择模式都是平行不周遍的,即不是所有的名词实现为单数和复数时都有清浊交替。因此从激发变音的条件看,无论是变体描写还是规则描写,简单程度相当。用变体描写,每个这样的语素有一个特殊的变体出现在复数前面,语素库要注明。用规则描写,每一个这样的语素要在语素库种做一个标注,指出它们需要选择特定的变音规则。两种方式增加的信息元数量是相同的。

当交替模式具有平行周遍性质时,选择规则描写方式就是必须的。比如,名词复数的变音和名词所有格词尾s、动词第三人称单数s以及分词和过去时标记ed可以统一起来。这些词尾有一个共同点,都含有真辅音(true consonant),它们的变音可以概括为以下两条规则:

1. 真辅音词尾在清音后读清音,浊音后读浊音。

2. 如果真辅音词尾前面是真辅音,发音部位都是舌尖音或舌叶音,并且在“塞”上相撞或“擦”上相撞(同时是持续的或同时是非持续的),必须在真辅音词尾前插入[i]后再使用规则一。(第2点也可以这样表述:如果真辅音词尾和前面同部位的辅音(舌尖和舌叶)同时是持续的或同时是非持续的,则必须在真辅音词尾前插入[i]后再使用规则一)

[tʃ]和[d]不是相同发音方法特征相遇,因为可以把[tʃ]分成两个音段,后一个音段是擦的特征,和[d]所带的“塞”特征不同,所以[tʃ]的“塞”不会延续到后面的d,过去分词或过去时语素{d}和前面的[tʃ]相遇不需要插入[i]。[tʃ]和[s]是相同特征相遇,因为[tʃ]所带的“擦”特征一直会传递到后面的s。

以上持续(continuant)和非持续(noncontinuant)主要根据Chomsky和Halle(1968, p299)的分类。不过Chomsky和

Halle 是把塞擦音 (the affricates) 作为非持续音。我们认为塞擦音是持续还是非持续, 和具体语言塞擦音的性质有关, 有的语言中塞擦音持续特征比较长, 比较明显, 有的比较短, 不明显。从下面读音看, 英语的塞擦音应该属于持续的:

That's right

这里 t's 的读音是塞擦, 来源于塞和擦的组合, 更一般地说, 来源于两个相同部位的组合。

把规则变音作为变体分析, 这种分析在两个语音形式差别不大的情况下, 容易接受和理解。但当两个形式变化很大, 比如像后面要分析的汉语的儿化, 这种处理就有问题了。考虑到汉语的儿化, 我们认为上面那种规则交替的情况可以只设立一个语素形式, 再在语音层面设立一条变音规则。

Harris (1942) 明确提出语素音位符号的概念, 比如 /najF/ (刀), 其中 /F/ 就是语素音位符号, 在复数 /-z/ 前是 /v/, 在其他情况下是 /f/, 类似的实例有 wife。Harris 还提出了语素音位公式, 即 /f/ 在复数 /-z/ 前被 /v/ 替代。Harris 的语素音位其实是音位的长度。Harris 承认, 这样的语素音位符号和语素音位公式在处理 ox 和 oxen 的关系时很麻烦。总之, 当时对语素音形还缺乏必要的认识, 语素音位理论实际上认为语素的形式由音位组成, 语素变体之间通常只在某个音位上的不同。

语素音位理论存在的最根本的问题是把语素音形在不同环境下的各种具体读音都看成是语素音位本身的一种变体, 不区分语素音形的规则变化和不规则变化。其实英语的复数、第三人称单数、单数所有格, 都有共同的音变规则。只要在语法层面给出了语法条件, 就可以在语音规则中实现变化 (见后面相关章节的讨论)。语素音位理论是建立在线性原则基础上的, 认为语素在任何情况下的组合都可以在线性方向上作表层切分。我们

所说的语素音形组合规则并不要求完全符合线性原则,而是认为语素组合分两步,一步是语素基本音形的线性连接,一步是施加语音规则,语音规则具有转换规则的性质。

Chomsky 和 Halle 也曾经使用过语素音位的概念,后来不再用这样的术语(见本书其他部分的讨论)。Chomsky 和 Halle (1968)的 formative 没有给出定义,但从用例看,相当于语素,其形式相当于语素音形。formative 的语音信息是存储在词库里面的。根据 Chomsky 和 Halle 的用例和行文可以看出,词库里储存的不是 formative 音形的变体,而是底层形式。formative 和 formative 的组合通过音系规则生成表层语音形式。所以,Chomsky 和 Halle 的理论跟 Bloomfield、Trubetzkoy 语素音位的观念有很大的区别,Chomsky 和 Halle 的理论是建立在规则基础上的,Bloomfield、Trubetzkoy 的语素音位理论是建立在分布基础上的。但是,由于 Chomsky 和 Halle 没有严格区分语音规则的周遍和不周遍问题,就不能有效获得真正的语素音形的底层形式,不能真正把语素音形变体和语音规则分开。

我们把语素音形(即语素基本音形)定义为不可以用周遍规则推导的语素的语音形式。语素音形和语素音位不同。语素音位并不区分规则和不规则的情况,凡是一个语素在不同的环境有不同的读音,都被看成语素变体。根据语素变体的思想,汉语有大量语素在儿化词中有变体,在上声前有变体。考虑“眼”的变体:

眼 1: -ian ²¹⁴	白眼 pai ³⁵ ian ²¹⁴
眼 2: ia ⁻²¹⁴	眼儿 iar ²¹⁴
眼 3: iam21-	眼眉 ian ²¹⁴ mei ³⁵ →iamei

语素音形基本上可以分成两种,成音节语素和不成音节语素。比如“梨”是一个音节,“葡萄”是两个音节,都是成音节语素。

音形。英语中 books 中的 s 是不成音节语素。汉语的儿化词是两个语素一个音节,所以儿化特征-r 也可以看出是不成音节的语素。

尽管人们也可以用交替分布来处理变调,比如说每个上声语素有三个音值:21、35、214,但从学习语言的情况看,学习者一般都是用规则而不是分布来处理连上问题。比如下面这个字:

锰 meŋ²¹⁴

告诉一个没有从理论上学习过连上变调规则的北京人,“锰水”怎么读,尽管他从来没有听说过这个组合,他仍然能够给出合理的变调结果。这说明北京人不是靠记住变体的方式来学习上声,因为在他没有听说过“锰”的 35 变体以前,不可能靠记忆发出正确的“锰水”读音。可见,上声变调作为一种周遍性的语音规则是母语人的语言知识,周遍规则处理的材料不是靠交替分布。一般地说,语素变音的变体描写方式不仅仅是不优化的问题,更重要的是描写的不充分性,正是后一个方面,决定了基本语素音形和规则描写的必要性。

当然,这取决于我们怎样理解变体和规则。结构语言学所说的变体有一部分也是有规则的,比如结构语言学处理上声变调要给出分布条件,这本身也是规则的体现。不过结构语言学并没有自觉地讨论音变规则,而是希望通过分布来控制规则。前面讨论过,规则可以区分出平行周遍规则和平行不周遍规则。平行不周遍规则可以理解成分布,比如:

削:ɕiao⁵⁵; ɕye⁵⁵

我们可以说“削”有两种分布,但不是所有的“ɕiao⁵⁵”都有“ɕye⁵⁵”这样的一种分布,也不是所有的“ɕye⁵⁵”都有“ɕiao⁵⁵”这样一种分

布。çiao 和 çye 的交替是不周遍的。总之,完全用分布来处理语音行为,不仅不够优化,而且不能反映语言知识,即不能解释语言单位和规则的生成性。

6.5 核心语音单位

语音单位首先可以从语音片断的长短或层次上来观察,由此可以分出音节、声母、韵母、韵、音位、区别特征等单位^①。这个视角只是我们要讨论的大前提。这里我们更关心的是从功能上看完全不同的几种语音单位,这就是核心语音单位、派生语音单位和表义语音单位。

我们所谓的核心语音单位必须满足两个条件:

1. 纯区别性:只能区别话语音形
2. 初始性:不能被推导出来

话语是指语素以及语素组合,话语音形是指语素以及语素组合的语音形式(简称语形)。在“广”([kuaŋ²¹⁴])这样的语素音形中,所包含的 k、u、a、ŋ、aŋ、uaŋ、214,都是核心语音单位,这类单位只区别语素音形,同时不能被推导出来。这类单位相互之间的组合关系和聚合关系是高度协和的(陈保亚 1988;1993,4.1.4),或者说有比较规则的格局(王洪君,1999,2.2.3)。从不同的层次看,核心语音单位可以是区别特征、音位(或音段)、韵母、声母、韵、声调等单位。

在语音单位层阶中,语素音形是一个重要的节点,语素音形内部各个层次的语音单位的行为规则和语素音形之间的行为规则有本质的区别。

^① 音步是临时组合,不算语音单位。

“二”语素音形中的-r 和“摠”语素音形中的-n, 都只在区别语素音形, 并没有直接作语素音形。-r 和-n 不同的是,-r 也能用来表达意义或作语素音形(陈保亚, 1988), 如[kar⁵¹](盖儿)中的-r, 除了能够把[kai⁵¹](盖)区别开, 还有体词和小称的意义。我们说分布在“二”音形中的-r 和“摠”音形中的-n, 只区别意义, 并且不能被推导出来, 是核心语音单位, 而分布在“盖儿”音形中的-r 不满足核心语音单位“只能区别话语音形”的条件, 因此不是核心语音单位。由于分布在“盖儿”音形中的-r 这样的语素音形长度正好和最短语音片断长度相等, 所以容易被看成是特殊的语音单位。汉语中的音节绝大多数都有表义的功能, 或者说直接作语素音形, 但由于不是最短语音片断, 通常不会被看成是特殊的语音单位。我们把能够直接用来表达意义的最小语音片断称为表义语音单位, [kuaŋ²¹¹](广)是表义语音单位, “盖儿”音形中的-r 是表义语音单位, “摠”中的-n 和“儿”中的-r 不是表义语音单位。容易看出, 这里的表义语音单位实际上就是语素音形, 我们有时候用表义语音单位这个术语, 只是为了理解和认识语音的方便。

汉语普通话音系中确实有核心语音单位-r, 但不是因为它出现在[kar⁵¹](盖儿)中而被确认为核心语音单位, 而是因为出现在[ər⁵¹](二)这一类字中。

汉语普通话的 o 和 ɤ 通常是不对立的, 但在下面两个字中对立:

o³⁵(哦)

ɤ³⁵(鹅)

“哦”是一个语素, 也是一个语气词, 并不满足只能区别话语音形这个条件, 所以这里的对立不是核心语音单位对立。在汉语普通话中, 音子 o 和 ɤ 是不对立的。王洪君(1999)认为语气词属于边际语音研究范围, 我们这里的分析能够证明这一点。再比较:

A⁵¹ (啊)ɤ⁵¹ (饿)

这里的“啊”也是一个语素,一个叹词,但是我们承认汉语普通话中有 A 和 ɤ 的核心对立,不是因为上面实例的对立,而是因为下面的对立:

k' A²¹¹ (卡)k' ɤ²¹¹ (渴)

即 A 和 ɤ 可以作为只区别话语音形的单位形成对立。

有一种语音单位尽管是纯区别性的,但可以通过规则推导出来。[ɕiə̃r⁵⁵] (星儿)中的ə̃和[ɕiər⁵⁵] (心儿)中的ə是对立的,但ə̃这样的派生单位在什么时候出现是可以通过核心语音单位和规则预测,如果一个语素音形的韵是 iŋ,儿化后就会派生出ə̃。这样的单位并不出现在单独的语素音形中。

这里有必要区分两种语音规则:语素音形内部语音规则和语素音形间语音规则。我们所说的核心语音单位是指语素音形内的单位,派生单位是指语素音形间使用语音规则后生成的单位。这种划分并不完全对应于早期生成音系学所说的深层语音单位和派生语音单位,早期生成音系学把语素音形内的音位变体也看成是派生单位(Chomsky, N. and Halle, M., 1968)。拿汉语的 a 音位来说,根据生成音系学的观点,可以有以下派生形式:

a → ɛ / i _ n

a → æ / y _ n

a → ɑ / _ ŋ, _ u

a → ʌ / _ * (后面没有韵尾)

a → a / 其他环境

在语素音形内部,生成音系学的这种观点和结构语言学的条件变体理论本质上是相同的,只有表述上的差别,即生成音系学在用上下文有关文法公式来表示条件变体。无论是结构语言学的条件变体还是生成音系学的上下文有关文法公式,这种派生观念带有研究者处理变体的态度。首先,是否是派生涉及归并音子(phone)的宽严。比如汉语拼音方案中*i*的派生形式有:

$$i \rightarrow \text{ɿ} / _ts, ts', s$$

$$i \rightarrow \text{ɿ} / _tʂ, tʂ', ʂ$$

$$i \rightarrow i / \text{其他条件}$$

但是,如果 ɿ 、 ɿ 、*i* 本身就被分成三个音位,*i* 就不存在派生的问题。两种分析态度从优选角度看是等价的。三分音位看起来增加了音位数目,但音位变体少了。合并的方式尽管音位数目少了,音位变体却增加了。当然也可能有人提出疑问,说上声字也可以分成两组,就不存在上声变调规则。但是这两类情况仍然有本质区别。在语素音形内部,规则的繁简程度是真正等价的。在语素组合中,用变体来处理上声字、儿化词根,就会使描写相当复杂。更重要的是,语素组合的规则确实存在的。比如,对于一个没有听说过的姓氏“米”,说汉语普通话的人能够正确说出“米老师”的声调。如果把上声语素处理成变体,在没有听说过“米老师”的情况,不可能断定“米老师”的声调。前面提到的“锰”也是这种情况。坚持语素音位原则的人可能拿出另一个理由,说每个上声语素都有一个出现在上声前的 35 变体。其实这是在用规则,不过没有用公式表述。

生成音系学派生的观念也受记音宽严的影响。如果记音宽一些,派生语音单位就会少一些:

$$a \rightarrow a / _n$$

$$a \rightarrow \alpha / _ \eta, _ u$$

$$a \rightarrow \Lambda / _ *$$

如果记音严一些,派生语音单位就会多一些:

$$a \rightarrow \epsilon / i _ n$$

$$a \rightarrow \text{æ} / y _ n$$

$$a \rightarrow \text{e} / u _ n$$

$$a \rightarrow \alpha / _ \eta, _ u$$

$$a \rightarrow \Lambda / _ *$$

$$a \rightarrow a / \text{其他环境}$$

因此,我们不把从语素音形内部切分出来的音看成是派生语音单位。我们限定的派生语音单位不取决于研究者记音的宽严和分类的宽严,只取决于核心语音单位和组合规则。比如无论我们把 $\text{ɪ}, \text{ɪ}, \text{i}$ 看成一个音位还是三个音位,通过核心语音单位和组合规则都只能在“影儿”中推导出派生语音单位 $\tilde{\text{a}}$ 。

关于语音规则类型的区分,前人已经做过一些研究,但都没有严格区分语素音形内部规则和语素组合的语音规则。Chomsky 和 Halle (1968, p13) 首先把音系规则分成两种,一种是用于短语的,一种是用于词的:

We will see that the phonological rules fall into two very different classes. Certain of these rules apply freely to phrases of any size, up to the level of the phonological phrase; others apply only to words.

Kiparsky(1982)在词汇音系学(Lexical phonology)中也曾把音系规则分成这样两种,一种是词的内部语音的规则,称为词

汇规则(lexical rules),另一种是跨越词界的后词汇规则(post-lexical rules),在词汇规则之后使用。Kiparsky(1982)进一步认为有三个由内向外的构词层面,比如:

	前置	后置
一层	illegal	dangerous
二层	nonillegal	dangerousness
三层		dangerousnesses

由此可以把附加语素分成三种:

	前置	后置
一层	il-	al,ous,ity,th
二层	non	hood,ness,er,ism,ist
三层		s,ed

可以看出,Kiparsky(1982)的词汇规则都属于语素音形组合成词的规则。词汇音系学本质上是把语素音形的组合分成两步来处理,词汇内和词汇外的,并没有把语素音形内部的规则和语素音形之间的规则区分开。实际上词汇语音规则和词汇外部语音规则的区别不是绝对的,前面已经提到过这一点,汉语普通话的上声变调是一个重要证据,两个字组合无论是词还短语,都可能发生连上变调。又比如德阳话后字变调:

词:腊肉、黄豆、胡豆、绿豆、白菜、油菜、麻布

词组:白肉、牛肉、鸭肉、驴肉、鱼肉、白布、蓝布、黑布

无论是词还是词组,阳平和去声组合都是 31+214→31-55。徐通锵(1997)主张的字本位,可以从这个角度加以认识。

赵元任(1968,1.3.9)提出过一种边际音位的概念,和我们上面提到的核心音位并不是对应的概念。赵元任的边际音位有

一部分相当于我们的核心语言单位中语音负担很少的音位,比如成都话的“合”和“核”的对立,有一部分包括语素组合中新出现的超出语素音形的对立,比如“歌儿”和“根儿”中的元音对立。后一种情况是可以用规则处理的。

20 世纪 70 年代后的生成音系学,尤其是词汇音系学,认识到由语法条件制约的语音规则与由纯语音条件制约的语音规则是不同的两类,它们在语言系统中的位置完全不同,这是 80 年代生成音系学研究的重要内容。Hooper(1976, p15; 1979, p107-108)区分了纯语音层面的规则(phonetic rules)和受词汇语法条件制约的规则(morphonemic rules)。中国学者张惠英(1979: 4)分析崇明方言的连读变调时,提出了“广用式”和“窄用式”两种不同的变调,“广用式”相当于纯语音层面的规则,“窄用式”相当于受词汇语法条件制约的规则。80 年代初期中国在这方面的研究比较活跃。

以上关于语音规则类型的各种区分和研究是极有价值的,但还不是语素内部语音规则和语素之间语音规则的区分。要展开语音调查的分析程序,后一种区分是必要的。语音最本质的行为差异就在于和意义的关系。当一个音形和意义有关系但没有和意义对应时,其分布是一种方式,比如汉语的 m。当一个音形和语音对应时,就产生了两个完全不同的结果:

1. 其分布就完全上升到符号层面,这时音形的分布是跟着语素分布走的。比如汉语的 r。
2. 出现变音现象。

正是根据这样的本质变化,我们可以把语音规则分成两种,一种是语素音形内部的规则,比如汉语普通话中 a 的三种读音;一种是语素组合形成的语音规则,如汉语上声变调、儿化韵、河南一带的 D 变韵、Z 变韵等。语素音形内部的语音规则是纯语音的,

并且是可逆的,比如,从汉语普通话拼音字母 a 根据一定的条件可以推出三个具体音值,而每个音值根据具体条件也可以推出音位。从根本上说这是一种处理技巧。语素之间的语音规则所产生的读音不一定有可逆性。从 35-214 推导不出是 35+214 还是 214+214。语素之间的语音规则可以进一步分出词内部的语音规则和短语的语音规则,语素之间的语音规则还可以分成纯语音的语音规则和受语法条件制约的语音规则,这些都是语素之间语音规则的进一步分类。前面说过,词内部语音规则和短语语音规则的区分并不是绝对的,这是因为词和短语的区分不是很清楚。就汉语的情况看,普通话儿化变韵是成词语音规则。两字组上声变调不分词和短语,是语素音形组合的语音规则。纯语音的语音规则和受语法条件制约的语音规则的区分也不是绝对的,这取决于我们所说的语法制约是什么含义。比如两个上声字的变调可以认为是纯语音的,但三个上声字的变调则要受到句法层次的影响,比如:

买/雨伞

雨伞/小

由于句法层次不同,这里的变调是有区别的:

买[雨伞]: 214+[214+214]→214+[35-214]

→214-35-214

[雨伞]小: [214+214]+214→[35-214]+214

→35-35-214

儿化变韵一般理解成纯语音的,但是否儿化则受到语法语义层面的影响。比如汉语普通话“婴儿”不儿化,“影儿”要儿化。当然也有些语素组合之间的语音规则是纯语音的,比如汉语普通话和很多语言有这样的规则:

[—鼻尾]+[唇—]→—m—唇

提取核心语音单位是音系研究的第一步工作。没有核心语音单位,就没法研究在语素组合中形成的语音组合规则和派生语音单位,也没有办法区分哪些是表达语音单位。因此,音系分析有三个重要任务:

1. 提取核心语音单位。
2. 确定语音组合规则和派生语音单位。
3. 找出受语法语义条件制约的各种表义语音单位和组合规则。

其实语素音形的重要性已经在 Chomsky 和 Halle(1968)中显示出来,因为那里的表层结构是 formative 构成的序列,根据我们对 Chomsky 和 Halle(1968)所用的材料的分析,formative 就是语素。在 formative 序列的基础上使用语音规则,就形成有具体音值的语音序列。至于怎样进一步分析语素音形,生成语法另有一套。

下面讨论提取核心语音单位的必要条件。

6.6 提取核心语音单位以语素基本音形的长度为必要条件

前面的分析实际上已经包含了这样一种观点:提取核心语音单位必须以语素基本音形的长度为必要条件。换句话说,在超出这个长度的语音序列中提取到的音段或音段组合就不一定是核心语音单位。下面我们来进一步论证这个问题。

先让我们再从长度上明确语素音形的含义。语言调查首先要获得大量独立出现的言语片断,下面几个层面的片断都可以

看成是独立出现的言语片断:

走、桌子、高

走了、第二

写一本书、谈两个问题

句子、文本也都可以看成是言语片断。在线性方向上最短的有意义的言语片断是语素。语素的语音形式可以称为语素音形。

要保证切分出来的语音单位是只能区别话语音形的单位,应该从语素音形入手,因为语素是最小的有意义的单位,从语素音形中切分出来的单位不再有意义,只起到区别语素音形的作用,这就能够保证了从语素音形中提取到的是纯区别单位或纯语音单位。如果从大于语素的片断来分析音系,比如从词、短语或更大的话语片断分析音系,就不能保证分析出来的最小语音片断一定是没有意义的片断。比如汉语[kar⁵¹](盖儿)中的-r是最短的语音片断;但不是纯语音单位。

容易看出,从语素音形描写音系的第一个理由是由自然语言的两层性决定的。前面的讨论已经从其他的角度涉及这个问题。自然语言的两层性是说自然语言有两种不同的单位,一种单位不仅能区别话语音形,还能表达意义,如语素以及由语素构成的一些言语片断,我们可以把这类单位称为表达单位或表义单位,另一种单位只能区别话语音形,我们把这类单位称为区别单位,也就是语音单位。表达单位的组合受语法、语义条件限制,区别单位的组合规则受纯语音条件限制,有时还会受到语法、语义的限制。音系研究的一个重要目标就是要找出区别单位及其在组合规则中的不同语音条件。Hockett、Martinet已经认识到了两层性,但并没有认识到语素基本音形内部片断和语素组合音形在分布上的根本区别,因此他们关于两层性的定义不够严密,这一点后面还会谈到。

当我们说核心语音单位的分析必须以语素基本音形为必要条件时,也就承认了音系分析不能独立于语法语义进行,这一观点和生成音系学的观点是一致的。生成音系学认为,音系分析不能独立于语法语义的信息展开,尤其需要词的概念,即词的界限、词根词和派生词等概念(Halle, 1959)。SPE 以后的生成音系学进一步认为需要把纯语音规则和跟语法相关的语音规则区别开。我们同意生成音系学关于语音单位的分析需要有语法层面的知识,我们的观点和生成音系学的区别在于,提取音系的核心语音单位和纯语音规则是以语素音形而不是词的音形为前提,也就是说,只有语素音形内部的单位和规则才是纯语音的。现在我们要追问:以词为测试环境的对比是否能有效的提取核心语音单位?更明确的追问是:以词为测试环境的对立是否能有效的提取只区别意义的单位?由于生成音系学不研究田野调查和分析程序,没有回答这样的问题。倒是结构语言学家 Hockett(1958, 2. 1, p17)认为用独词句来确定音位是一种方便的做法, Bloomfield(1933, 5)实际上也是这么做的。下面我们通过汉语实例来分析以词为测试环境所遇到的困难。

困难首先在于这种办法不能有效区分核心语音单位和表义语音单位。考虑下面三个音形的对比:

kar⁵⁵ 杆儿

kuan⁵⁵ 干

kaŋ⁵⁵ 钢

根据对比原则,可以断定-r 和-n、-ŋ 是 3 个不同的音位,都可以作为韵尾出现。但是-r 的分布和-n、-ŋ 的分布很不相同(陈保亚, 1988), -n、-ŋ 的组合很有规则,只分布在几个主要的元音后:

	a-	ə-	i-	u-
-n	an	ən	in	□
-ŋ	aŋ	əŋ	iŋ	uŋ

和-n、-ŋ相比,-r的组合行为似乎很没有规则,能分布在很多元音或元音组合后:

	a-	ə-	i-	u-	o-	ε-	ɤ-	ou-	au-	ā-	ǎ-	ĩ-	ũ-
-n	an	ən	in	□	□	□	□	□	□	□	□	□	□
-ŋ	aŋ	əŋ	iŋ	uŋ	□	□	□	□	□	□	□	□	□
-r	ar	ər	□	ur	or	εr	ɤr	our	aur	ār	ǎr	ĩr	ũr

按照结构语言学的方法提取出来的-r和-n、-ŋ都是音位,它们在分布上却有如此大的差异,原因在于没有先提取语素,因此对比的不全是语素音形。在上面的实例中,“杆儿”是两个语素构成的词,“干、钢”各是一个语素,也各是一个词。由于不是在语素层面上对比,因此对比出来的单位可能是纯区别性的核心语音单位,比如-n、-ŋ,也可能既是纯区别性的核心语音单位,比如“儿”中的-r,也是表义语音单位,比如“杆儿”中的-r。核心语音单位是只区别语素音形的单位,北京话中-n、-ŋ是这样的单位,它们的分布和语法语义没有关系,从根本上说和语素组合没有关系。-r不纯粹是这样的单位,“杆儿”中的-r的分布和语法语义条件有关系,它和前面的语素组合后形成有小称功能的词,而且绝大部分是名词,所以可以广泛分布在很多语素音形后面。

从以上-n、-ŋ和-r的分布差异立即可以看出,核心语音单位和非核心语音单位在分布上有根本的差异,过去由于没有充分认识到这一点,核心语音单位的价值也没有被充分认识到。Hockett(1958)曾经讨论过两层性(duality),即音系系统(phonological system)和语法系统(grammatical system)。Martinet

(1960)给出了双重分节(La double articulation)的概念,包括第一分节(La première articulation)和第二分节(La deuxième articulation),第一分节的最小单位是 monème,相当于语素(morpheme),第二分节的单位是 phonème,相当于音位(phoneme)。Hockett 和 Martinet 都认识到了符号的两层性,也认识到了音系层可以为语法层提供各种符号形式,从几十个音位,可以产生成千的音节,为语素提供了足够多的音形。但他们区分的音系系统不仅仅包括核心语音单位,也包括非核心语音单位,所以核心语音单位的性质并没有被充分认识到。如果把这两种单位混在一起,就无法认识核心语音单位的分布规律。如果把核心语音单位和非核心语音单位区别开,与此相应的语音规则也有两种,一种是语素音形内部的,是纯语音的,一种是语素组合的语音规则,可能是纯语音的,也可能是和语法语义有关系的。语素音形内部的语音单位的分布和规则不受语法层面因素的影响。

确定纯区别性的核心语音单位应该从最小的有意义的片断的形式入手,当这样的形式再往下切分时,它的任何一部分都不再表达意义,根据这样一些最小有意义片断的形式切分出来的语音片断才能保证是纯区别性的核心语音单位。而这样的最小的有意义的片断就是语素。如果语素的语音形式本身已经不能再切分,比如汉语表示小称的-r,英语表示复数的-s,这样的片断不仅是语素,也是音位,它们的分布条件既涉及语法,比如儿化词中的-r,也涉及语音,比如“耳”中的-r(陈保亚,1988)。如果从词入手来提取音位,由于词可以是两个或两个以上的语素的组合,从词的音形中切分出来的片断就有可能是有意义的单位,如 kar⁵¹(盖儿)中的-r。

汉语方言调查的主要方式长期以来先调查单字音,再描写变音,这一方法通常是有效性的,这是因为汉语的语素音形基本上是单字音。生成音系学 SPE 模型主张在词库中存放 formative 深层形式,通过音系规则获得表层语音形式,也是从语素音

形出发的一种体现。不过这两种方法的提倡者都没有从根本上论证为什么一定要从 formative 音形出发,没有论证语素音形内部的语音和语素音形组合后出现的语音的根本区别。

我们曾经认为汉语普通话儿化导致音系不协和(陈保亚, 1988),其实问题的实质更本质上还是在于基本音形和变音的差异。在基本音形中,ər 只和零声母组合,ər 在基本音形构成的系统中是一个高度不协和的音。可见基本音的系统不一定总是协和的。儿化音节本身,从韵母和声母的组合看,倒是相当协和的。

是否可以考虑从单纯词入手来提取核心语音单位? 单纯词是只含有一个语素的词,根据这样的定义,儿化词就是合成词而不是单纯词,在分析程序中不用儿化词这样的合成词形式,只用“干、钢”这样的单纯词形式,似乎就可以避免从儿化词的语音形式中提取核心语音单位所带来的困难。但是,判定单纯词和合成词必须假定语素已经提取,所以用单纯词提取核心语音单位本质上还是从语素音形入手。另外,从单纯词入手提取核心语音单位可能是不充分的,单纯词都是可以单说的语素,而有些核心语音单位可能只存在于黏着语素中,比如“佛”([fo³⁵]),这个音形所辖的其他语素“拂、埤、伏”等也都是黏着语素。这些语素的音形是不能在音系中忽略的,否则所得到的语素音形不充分,所得到的小于语素音形的片断也不充分。

这里我们只是论证了从词出发不能有效提取到纯区别性核心语音单位。从词入手在田野调查中还有一个困难,就是必须假定我们已经明确提取了一个语言中所有的词。以 Bloomfield 为代表的早期结构语言学在分析音位时一般假定词的存在,生成音系学也假定词的存在,事实上判定词是相当困难的(陈保亚, 1999, 6)。尤其是在汉藏语言中,提取词比提取语素要困难得多,用词来提取区别性单位的条件更不成熟。

从上面的分析立刻可以看出,通过音节提取核心语音单位也不具有充分性,会把核心语音单位、派生语音单位和表达语音

单位混在一起,因为“玩儿、心儿、肚儿、盖儿”等儿化片断都是包含两个语素的单音节形式,这样的音节在汉语中大量存在,汉语方言中的 Z 变韵、D 变韵也属于这种情况。另外,在很多语言中,语素音形的切分和音节的切分不总是一致的,音节本身的提取已经遇到了困难,比如英语的 working:

语素音形	wə:k	ɪŋ
音节	wə:k	kɪŋ

k 在这里跨越两个音节,严格地说更可能属于后一个音节。又比如 actor:

语素音形	ækt	ər
音节	æk	tər

这里音节和语素的界限不在同一个接点上。所以,尽管从诗歌对仗的证据看,音节是说话人容易感知出来的语音片断,尤其是在汉语中,每个音节都有一个声调,音节容易通过声调来识别,但音节不宜作为提取核心语音单位的起点。

Harris(1951)还使用过另一种方法,启用比词更长的话语为测试环境提取音位,这是后结构语言学发现程序的思路,目的是在提取音系单位时只从调查到的话语片断入手,完全不考虑语素、词等语法语义层面的单位,生成音系学已经从规则描写的简单性上分析了这种方法的弱点,此处不再重复。从核心语音单位的提取看,这种方法所面临的困难和从词入手面临的困难是相同的。观察下面的实例:

She seems ill [ʃi: si:mz il]
She seemed ill [ʃi: si:md il]

[-z]和[-d]两个形式是这两个话语片断的不同之处。这种办法

仍然不能说明对比出来的[-z]和[-d]是否是纯区别性核心语音单位,或者说[-z]和[-d]是否是纯音位并不能由此得到证明。我们认为,[-z]和[-d]的音位对立应该是通过类似下面语素音形证明的:

zoo [zu:] 动物园

do [du:] 做

概括地说,从话语提取音位有两个问题:

1. 提取不充分,因为在语素层面对立的音在句子层面不一定能找到对立实例
2. 有些对立的音子并不是基本音形。

其实核心语音单位如果能从话语的对比中提取出来,也能够从语素音形中提取出来。观察下面实例:

mai21p'an³⁵tsɿ 买盘子

mai21p'ən³⁵tsɿ 买盆子

如果认为上面两个话语片断的区别是 a 和 ə 决定的,这里的 a 和 ə 也直接可以从“盘、盆”两个语素音形的对立中提取出来。一般地说,如果两个核心音位能够从言语片断的对比或词的对比中提取出来,也应该从语素对比中提取出来。反过来就不成立,从比语素长的词或言语片断中提取的片断可能是纯区别性核心语音单位,也可能不是纯区别性核心语音单位,而从语素音形对比中提取到的总是纯区别性核心语音单位。当语素音形的长度只等于一个音位时,比如汉语“杆儿”中的-r,英语的复数语素-s,可以通过其他语素音形来证明该片断是否是核心语音

单位。

纯区别性的核心语音单位的分布以语音为条件,有意义的单位的分布通常不受纯语音条件的限制。自由语素肯定不受任何语音条件的限制,黏着语素由于分布环境是有限的,受到一定限制,但其语音形式的分布从纯区别单位的角度看是不规则的。

音系分析首先要找到可以只区别语素音形的核心语音单位,找出这些单位的分布规律。在没有首先提取到语素的前提下,从话语、词等入手不能达到这一目的。我们说从话语、词等入手不能达到这一目的,是指不能直接达到这一目的。有人可能会考虑修正结构语言学的方法,先不提取语素,直接从话语或词入手,先提取对立的音子,再区分哪些是只区别意义的音子,哪些是既可以区别意义,也可以直接作语素音形的音子。具体地说,直接从词形或言语片断中先提取到-n、-ŋ、-r 这样一些音段,最后再确定-n、-ŋ 是只区别意义的音段,-r 在“二、而、尔、耳”等语形中是只区别意义的音段,在“门儿、心儿、味儿”等语形中既区别意义,也表达意义,然后再说明-n、-ŋ 在音系层面的分布是规则的,-r 在音系层面的分布是不规则的。但是,这样做暗中仍然需要在提取语素的前提下进行,即我们必须知道-n、-ŋ 不是语素,“门儿[mər³⁵]”中的-r 是语素,“二[ər⁵¹]”中的-r 不是。

核心语音单位的第二个条件要求核心语音单位不能是派生的,这要求我们有办法来区分初始和派生两种现象。一般地说,派生语音单位最本质的特点是产生于语素组合中,比如ā的分布和“儿化”有关系,“的”读音的不同高低和前面语素音形的声调有关系。在平舌韵中,中元音是不对立的,都是互补的,但在儿化词中,中元音产生了对立。Hartman(1944)、Hockett(1947)、赵元任《汉语口语语法》(1968,1.3)都观察到了这一点。我们可以把带中元音的音节儿化以后产生的对立归纳成下面几种:

儿化类型		实例 ^①	
ən、əi+r	ɤ+r	根儿[kər]	歌儿[kɤr]
		泪儿[lər]	乐儿[lɤr]
ɿ+r	ɤ+r	纸儿[tʂər]	褶儿[tʂɤr]
ən、əi+r	o+r	魂儿[xuər]	活儿[xuor]
		鬼儿[kuər]	果儿[kuor]
		盆儿[p'ər]	婆儿[p'or]
		轮儿[luər]	锣儿[luor]
		滚儿[kuər]	果儿[kuor]
		姨儿[iər]	爷儿[ier]
i+r	iɛ+r	鸡儿[tɕər]	街儿[tɕer]
		须儿[ɕyər]	靴儿[ɕyer]
y+r	yɛ+r		

这些对立的儿化词都体现为[ə]和[ɛ、ɤ、o]的对立。但是这种对立的性质和核心语音单位对立的性质是完全不同的,这种对立只是发生在语素和语素相组合的规则变化中,不出现在单个的语素对立中。这种对立完全可以通过语音规则推导出来。比如,[ə]和[ɛ]的对立完全可以由[i]韵母和[iɛ]韵母的对立在儿化条件下推导出来:

$$i + \text{ər} \rightarrow i\text{ər}$$

$$i\text{ɛ} + \text{ər} \rightarrow i\text{ɛr}$$

儿化词除了导致上述元音的对立,还导致了鼻化元音的产生,并形成鼻化元音和非鼻化元音的对立:

^① 赵元任(1968,p33)认为“根儿”跟“歌儿”的区别,“鸡”跟“街儿”的区别,正在很快地消失之中。

kār 缸儿	kar 杆儿
ɕiər ⁵⁵ 星儿	ɕiər ⁵⁵ 心儿

鼻化元音的出现以及口音、鼻音对立,和核心语音单位及其对立性质是不同的,都是从词音形对比的结果。

派生对立不仅仅限于语素组合成词的过程中,根本上是语素组合的结果。汉语的变调就不同于语素成词的过程中。比如:

tsuo ⁵¹ tɕ'uan ³⁵ pian ⁵⁵	坐船边
tsuo ⁵¹ tɕ'uɑŋ ³⁵ pian ⁵⁵	坐床边

本来 a 和 α 不对立,由于下面两条规则:

$$\begin{aligned} _n + p_ &\rightarrow _mp_ \\ _ŋ + p_ &\rightarrow _mp_ \end{aligned}$$

导致上面两个句子产生以下变化:

$$\begin{aligned} \text{tsuo}^{51} \text{tɕ'uan}^{35} \text{pian}^{55} &\rightarrow \text{tsuo}^{51} \text{tɕ'uam}^{35} \text{pian}^{55} \\ \text{tsuo}^{51} \text{tɕ'uɑŋ}^{35} \text{pian}^{55} &\rightarrow \text{tsuo}^{51} \text{tɕ'uɑm}^{35} \text{pian}^{55} \end{aligned}$$

现在在组合中产生了对立。

Hockett(1947)以及后来的好些学者在描写北京音系的时候,把-r 和-n、-ŋ 形式作相同对待,把[ə]和[ɛ、ɤ、o]作相同对待,不能解释-r 的分布和-n、-ŋ 的分布的差异,不能解释-n、-ŋ、-r 对立的特殊性,不能解释[ə]和[ɛ]这类对立的特殊环境,也不能解释鼻化元音的产生以及口音、鼻音对立形成的特殊机制,最终原因就在于没有从语素音形入手来提取核心语音单位。由此可见,用于提取纯区别性核心语音单位的片断不应该大于语素音

形。语素基本音形是纯区别性核心语音单位描写的起点。

其实语流中的变音是相当复杂的,我们现在研究得很不够。赵元任(1968,p37—38)有以下例子:

$t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{ən}^{35}\text{mə}$	这是什么?
$t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{y}^{35}\text{mə}$	这是蛇吗?

赵元任认为这两句读音很相似。按照赵元任的理解,产生了下面的语流音变:

$t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{ən}^{35}\text{mə} \rightarrow t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{ym}^{35}\text{mə}$	这是什么?
$t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{y}^{35}\text{mə} \rightarrow t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{ym}^{35}\text{mə}$	这是蛇吗?

实际上更多的情况下是产生了一种对立。变化规则如下:

$t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{ən}^{35}\text{mə} \rightarrow t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{ə}^{35}\text{m}^{35}\text{mə}$	这是什么?
$t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{y}^{35}\text{mə} \rightarrow t\text{ɕ}\text{y}^{51}\text{ɕ}\text{l}^{51}\text{ɕ}\text{ym}^{35}\text{mə}$	这是蛇吗?

即形成了[ə]和[y]的对立或[ə]和[ə]的。

从以上分析可以看出,核心语音单位和派生语音单位最根本的区别之一是,核心语音单位不能通过语素组合推导出来,派生语音单位能够通过语素组合推导出来。从语素基本音形中切分出来的片断不仅是纯区别性语音单位,而且是非派生的初始语音单位,即这些单位不可以通过规则推导出来。

6.7 提取核心语音单位的语境不能短于语素音形

如果核心语音单位的提取要求分析的语料不能长于语素音形,那么是否可以短于语素音形?有些音节是短于语素音形的,

比如[p'u³⁵t'au⁰](葡萄)中的音节。前面我们已经讨论过从音节入手不能保证得到的片断是纯区别单位和非派生单位。还有一些短于语素音形的单位,比如说韵,可以直接在诗歌的押韵中提取出来。有些双声叠韵词,也可能从中把声母和韵母提取出来,那么核心语音单位的提取是否可以建立在这样一些语音片断上?下面的分析证明不能。

核心语音单位一个重要属性是对立。结构语言学重视音子的对立,提出了音位的概念。生成音系学曾主张抛弃对立、音位这样一些概念,用音段或语音特征这样一些概念来描写音系。但是,对立的概念是避免不了的,拿音段来说,生成音系学不会把汉语中两个对立的音子[ts]和[tɕ]看成同一个音段,音段概念本身已经隐含了对立的概念或音位的概念。另外,任何做方言调查和民族语言调查的人都知道,对立和音位的概念也是绕不过去的。区别特征也是和对立有关系的。因此提取音位、区别性特征等对立单位是充分提取核心语音单位的主要任务。问题是怎样提取这些单位?结构语言学的发现程序是先提取音位,然后再提取语素,提取音位时提出了对立原则和互补相似原则,这一操作方法假定音系的分析可以独立于语素,可以直接从音位或区别特征入手。实际操作中这种分析手续会遇到困难。根据结构语言学提取音位的互补相似原则,下面的两个辅音是互补的:

tɕi⁵¹ 继

tsɿ⁵¹ 自

如果把它们看成相似,可以归并成一个音位j。同样根据互补相似原则,上面的两个元音也是互补的,如果把它们看成相似,也可以归并成一个音位i。仅仅根据互补相似原则,辅音和元音都有理由归并成同一个音位,于是就有:

ji⁵¹ 继ji⁵¹ 自

但实际上没有人会这样做,因为这两个语素音形不同,发音人很容易判断这一点。可见在分析音位时,首先要知道语素音形是否对立,语素音形的对立是根本的对立,是音位分析的起点。生成音系学不可能否认以上两个语素音形的对立。当音子不对立但语素音形对立时,必须依照语素音形的对立处理音位。其实在实际工作中,结构语言学处理元辅音的对立时,并没有把元辅音彼此隔离开来观察对立,只是没有自觉意识到并提出语素音形的对立是最根本的对立。

以上关系说明即使 tɕ 和 ts 不对立, i 和 ɿ 不对立, tɕi⁵¹ 和 tsɿ⁵¹ 也可以对立,因此在音系中,低层面单位音子的对立情况并不能完全推导出高层面单位语素音形对立的情况。汉语中 tɕ 和 ts 不对立, i 和 ɿ 不对立,可以进一步概括成舌面和舌尖前也不对立,因此语音特征不对立也并不等于语素音形不对立。语素音形的对立是音系中最初始的对立。

与对立密切相关的一个问题是对立片断的长度。一般认为汉语普通话的 -n 和 -ŋ 有对立,那是以归并韵腹或不考虑韵腹的差异为前提。比如:

船	tɕ'uan ³⁵
床	tɕ'uaŋ ³⁵
心	ɕin ⁵⁵
星	ɕiŋ ⁵⁵
笨	pən ⁵¹
蚌	pəŋ ⁵¹

但更严格的记音显示,以上的对立也可以是韵的对立而不只是

辅音韵尾 n 和 ŋ 的对立,也就是说,对立的负担可能并不仅仅落在韵尾上:

船	tɕ'uan ³⁵
床	tɕ'uaŋ ³⁵

心	ɕin ⁵⁵
星	ɕiŋ ⁵⁵

笨	pɔn ⁵¹
蚌	pɔŋ ⁵¹

这一点还可以从连读过程中-n 和-ŋ 的中和进一步证实。根据目前对汉语普通话的音段分析,如果不考虑 a 和 ɑ 的区别,前面提到的两句话在连读过程中由于唇音的影响会变成同音的句子:

tsuo⁵¹ tɕ'uan³⁵ pian⁵⁵ → tsuo⁵¹ tɕ'uam³⁵ pian⁵⁵ 坐船边

tsuo⁵¹ tɕ'uaŋ³⁵ pian⁵⁵ → tsuo⁵¹ tɕ'uam³⁵ pian⁵⁵ 坐床边

这两个片断中,“船”和“床”两个语素音形中的-n 和-ŋ 都变成了[m],如果“船”和“床”读音的区别仅仅在韵尾,那么在连读过程中这两个言语片断应该完全同音,但实际上“船”和“床”所在的两个话语片断读音仍然有区别,所以这里的对立负担不仅仅落在-n 和-ŋ,而可能是[tɕ'uan]和[tɕ'uaŋ]整个语素音形中更长的语音片断的对立,即韵的对立。从另一个角度看,也可以说我们在记音和归纳音位的时候没有把前 a([tɕ'uan]中的 a)和后 a([tɕ'uaŋ]中的 a)分开,如果分开了,这两个话语在语音上的区别就不仅仅是 n 和 ŋ 的区别,连读变音后仍然保留区别就可以得到解释。实际上的读音和变音规则前面已经给出:

tsuo⁵¹ tɕʰuan³⁵ pian⁵⁵ → tsuo⁵¹ tɕʰuam³⁵ pian⁵⁵ 坐船边

tsuo⁵¹ tɕʰuaŋ³⁵ pian⁵⁵ → tsuo⁵¹ tɕʰuəm³⁵ pian⁵⁵ 坐床边

即“船”和“床”中的 a 和 α 在起对立的作用。再比如,按照目前对汉语普通话的音段分析,有如下连读变音规则(忽略儿化):

ɕin⁵⁵ mei³⁵ le → ɕim⁵⁵ mei³⁵ le 心没了

ɕiŋ⁵⁵ mei³⁵ le → ɕim⁵⁵ mei³⁵ le 星没了

连读变音的结果应该是两句话变成相同了。实际上,这两句话是有区别的,遵循下面的变音规则:

ɕin⁵⁵ mei³⁵ le → ɕim⁵⁵ mei³⁵ le 心没了

ɕiŋ⁵⁵ mei³⁵ le → ɕim⁵⁵ mei³⁵ le 星没了

即“心”和“星”中的元音在起对立作用。一般地说,两个语素音形的对立是否最终可以归约为一个音子的对立而不是两个或两个以上音子串的对立,在某些情况下是不清楚的,需要深入研究,尤其是在有些声调语言或方言中,某些对立是声调造成的还是声母造成的,抑或是声调和声母共同造成的,并不是很清楚。在汉语的声调阴阳取代声母清浊的演变过程中,应该有一个阶段是清浊和阴阳共同起区别作用。

其实现代语音学已经注意到,说话时,舌、唇、喉的运动是相互作用的,并不是像音段理论那样认为是同时开始和同时结束。Goldsmith(1979)的自主音段音系学(Autosegmental Phonology),更进一步说明各种语音特征在线性方向的传递距离是不完全一致的。Goldsmith(1979, p202)主张多线性音系分析(multilinear phonological analysis),不同的特征被置于不同的层面(tiers),不同层面的特征并不总是一一映射的(one-to-one map-

ping)。至于映射的情况最终是什么,往往需要深入分析才能得到。很多问题,包括普通话的语音,现在我们都还不是很清楚,但无论是哪一种情况,语素音形的对立是没有疑问的,语素音形应该首先确定下来。

6.8 语素音形的一致性和独立语素音形的确定

前面我们把满足纯区别性和初始性的语音单位看做是核心语音单位,从语素基本音形中切分出来的语音单位正好满足这两个条件。所以我们可以给核心语音单位一个操作上的定义:核心语音单位是指从语素基本音形中切分出来的各种语音单位(不包括语素音形本身)。

既然提取核心语音单位的语音片断既不能长于语素基本音形,也不能短于语素基本音形,而核心语音单位的提取又是音系分析的基础,那么提取语素和得到基本语素音形就是音系描写的起点。由于基本语素音形在不同的语境中的语音表现不是唯一的,在建立语音规则以前,语素基本音形并不能确定,因此提取语素基本音形也必须同时考虑语音规则。

再回到结构语言学的发现程序上来。我们说结构语言学从语音到语素的发现程序存在问题,分析应该从语素音形到语音单位,这就涉及如何独立提取语素的问题。前面已经讨论过, Saussure(1912)最先提出提取单位的基本原则是对比,后来 Bloomfield(1926, 1933)以及结构语言学家深入讨论了对比。如果给出一定的限制(陈保亚, 2006, 2),用对比原则从记录下来的言语片断中提取语素基本上是可行的,从语子具体归纳语素的时候会有一些分歧,但语素音形的一致性支持从语素音形提取核心语音单位的可行性。语素音形的一致性是指语素音形的同一性在不同的调查者之间基本上是一致的,容易得到实证,容易在研究者之间达成共识。语素音形的一致性特点保证了从语

素音形提取核心语音单位的有效性。语素音形可以被切分成更短的语音片断,最短的语音片断是音子(phone)或音段。和语素音形比较,语音片断和音子的一致性要弱一些。以语素音形和音子的比较为例,不同调查者记音的宽严标准不同,受母语干扰的程度不同,听音的准确程度不同,所记录出来的音子的数量也不同。仅以汉语普通话“前、全”两个语素音形的元音记音情况为例,下面的情况都有可能:

	前	全
调查者 1	tɕ'ien ³⁵	tɕ'yen ³⁵
调查者 2	tɕ'ien ³⁵	tɕ'yən ³⁵
调查者 3	tɕ'ien ³⁵	tɕ'yæn ³⁵
调查者 4	tɕ'ien ³⁵	tɕ'yen ³⁵
调查者 5	tɕ'iaen ³⁵	tɕ'yæn ³⁵
调查者 6	tɕ'iaen ³⁵	tɕ'yen ³⁵
调查者 7	tɕ'ien ³⁵	tɕ'yən ³⁵
调查者 8	tɕ'ien ³⁵	tɕ'yen ³⁵
调查者 9	tɕ'ien ³⁵	tɕ'yen ³⁵

不同的调查者可以有不同的记音,主要元音在音值和数量上都有差别,或许还有更复杂的情况,但无论是哪一种情况,无论对哪一个调查者,这里都只有两个基本语素音形,这对所有调查者来说都是一致的。

在把语素音形切分成更短的语音片断时,不同的研究者处理方式也可以不一样。考虑下面三个语素音形的音段情况:

汉语拼音	zhan ⁵⁵	chan ⁵⁵	shan ⁵⁵
语素	粘	掺	删
方案 1 的构造	tʂ+a+n	tʂ'+a+n	ʂ+a+n
方案 2 的构造	ts+r+a+n	ts'+r+a+n	s+r+a+n
方案 3 的构造	t+s+r+a+n	t'+s+r+a+n	s+r+a+n

方案 1 是汉语拼音方案切分音段的方式。方案 2 把声母分析成舌尖前音和 r 的组合,这是 Hartman(1944)分析北京话时采取的方式。方案 3 是在方案 2 的基础上进一步把塞擦音分析成塞音和擦音。把塞擦音看成塞音和擦音的组合,可以从英语的合音得到支持:

What is this→ What's this

尽管这 3 种方式切分出来的音子数目不一样,但基本语素音形数目是不变的,对所有分析者来说都是一致的。

在把音子归纳为音位的时候,由于涉及互补相似原则,即使在记音完全相同和线性切分完全相同的情况下,不同的研究者归纳出的音位也不一定相同,音位变体也不一致。考虑下面 3 个语素的音质音位:

子	几	纸	音位归纳方案
tsɿ	tɕi	tʂɿ	记音形式
tsɿ	tɕi	tʂɿ	音位归纳方案 1:6 个音质音位
tsi	tɕi	tʂɿ	音位归纳方案 2:5 个音质音位
tsɿ	tɕi	tʂi	音位归纳方案 3:5 个音质音位
tsɿ	tɕi	tʂi	音位归纳方案 4:5 个音质音位
tsi	tɕi	tʂi	音位归纳方案 5:5 个音质音位
tsi	tɕi	tʂi	音位归纳方案 6:4 个音质音位

以上只考虑了三个语素音形中音子的归并。如果考虑更多的语素音形,情况更复杂,但无论哪一种处理方式,以上语素基本音形的数目都是一样的,只有 3 个,这对所有的调查者都是一致的。

当然,语素基本音形也有不同的表现或不同的变体,但是,只要把满足平行周遍条件的变音用规则来处理,剩下的语素音形的变体数量对不同的调查者来说也是相当一致的。下面给出

一个确定汉语普通话语素基本音形表的例子,然后逐一解释确定语素基本音形表的理由,这种做法代表了我们在汉藏语言调查中的一般做法。

[汉语普通话语素基本音形表]

语素基本音形	语素 1	语素 2	语素 3	
pan ⁵⁵	班	般	搬	
mou ³⁵	谋	牟	眸	
u ²¹⁴	五	梧	武	
u ³⁵	无	吴	芜	
ʂui ³⁵	谁			
ʂəi ³⁵	谁			
i ⁵⁵	一	衣	依	
i ⁵¹	一	意	义	
ia ⁵¹	亚	压	娅	
ia ²¹⁴	亚	雅	哑	
-r(-ər)	-儿			
lə	了			
la ⁵⁵ tɕei ⁵⁵	垃圾			
suo ²¹⁴	锁	所	琐	
ər ⁵¹	二	贰		
liɑŋ ²¹⁴	两	辆		

“班、般、搬”这样的语素音形都只有一个表现形式,通过对比直接存放在语素音形表中,同音的存放在同一个语素音形下。

语素音形[mou³⁵](谋)和[pan⁵⁵]的一个区别是,[mou³⁵]这样的音节从来不单独出现。

“五”有 u²¹⁴ 和 u³⁵ 两个表现形式(或记录形式),音高不同,由于 u³⁵ 中的 35 是可以通过 214 后面跟 214 推导出来,所以在 u³⁵ 记录中不再出现语素“五”。为什么不把“五”的表现形式 u³⁵

作为基本形式存放在语素音形表中呢？也就是说，如果我们假定 u^{35} 是基本形式， u^{214} 就是可以推导的，规则就是不出现在 214 前。但是这种处理不能解释为什么在有的情况下“五”和“无”有对立。从另一个角度看，如果 u^{35} 记录下只含有“无、吴”这样一些语素，不包括“五、捂”这样一些语素， u^{35} 栏中这些语素音形的活动方式就是一样的，但如果是“无、吴、五、捂”这样一些语素都在一起，语素基本音形的活动方式就不一样了。一般的说，两个语素音形只要有一处的分布读音是对立的，就是不同的语素音形，这个对立的读音通常应该是基本形式，才可能有效使用语音规则。

“谁”有 $[ɕui^{35}]$ 和 $[ɕəi^{35}]$ 两个音形。这两个音形可以和其他语素音形形成对立，肯定要收入对立语素音形表。因为音系分析的本质是要找出对立系统。当然音系规则或词库中还应该标出出现条件。

“一”和“亚”各自都有两个表现形式，都出现在两个记录中，但有区别。对“一”来说，在语音规则部分还要设立一条语音规则，说明“一”的出现条件，而“亚”不需要这样的语音规则。尽管语音规则部分有“一”的读音规则，但语素音形表中仍然要在两个语素音形下列出语素“一”，因为不是所有的 $[i^{55}]$ 都可以通过规则推导出 $[i^{51}]$ ，也不是所有的 $[i^{51}]$ 可以通过规则推导得出 $[i^{55}]$ 。在这一点上， $[i^{55}]$ （一）的变调性质和 $[u^{214}]$ （五）的变调性质是不一样的。

“一”和“衣”单说的时候，读音相同，但“一”在其他声调前和“衣”的声调不同，如果我们根据对立把“一”的 $[i^{51}]$ 看成基本的，又会和“意”的声调一样。比较合理的办法是把分布广的作为基本形式，这就是 $[i^{55}]$ 。

一般地说，如果一个语素的表现形式和其他语素的表现形式不同，或者说有对立，并且不能通过规则推导出来，都应该在语素基本音形表中占据一个记录。

后附的黏着语素“了”，有 4 个不同的高低形式，都是可以通过规则判定的：

阴平 + 轻音	阳平 + 轻音	上声 + 轻音	去声 + 轻音
k'u ⁵⁵ lə ² 哭了	lai ³⁵ lə ³ 来了	tsou ²¹⁴ lə ⁴ 走了	tsuo ⁵¹ lə ¹ 做了

这样的语素音形在记录表中可以不标出音调的高低，但在规则库中要给出规则。

儿化的-r，也是一个语素音形，规则表中要给出儿化规则。

“垃圾”尽管是两个音节，但只是一个语素音形，只需要一个记录。

提取语素时由于涉及意义的同一性问题，不同的学者会有分歧，比如“锁门”的“锁”和“上锁”的“锁”，是一个语素还是两个不同的语素。在词平面也存在是一个词还是两个不同的词的问题。尽管有分歧，但语素基本音形只是一个，这对所有的调查者都是一致的。“二”和“两”，是一个语素还是两个不同的语素，也有分歧，但语素基本音形是两个，这对所有的调查者也都是一致的。

正是语素音形的一致性特性，使语素音形成为提取核心语音单位的基本前提。在汉语方言或汉藏语的民族语言调查中，如果提取了语素基本音形表，就可以得到各层面的核心语音单位，得到由这些核心语音单位构成的各种规则。但是，如果只给出音位及其可能的组合规则，还不能保证能够得到全部语素基本音形及其活动规则。

现在回到判定语素基本音形根本原则。在可观察语料中，一个语素的表现形式可以是不同的，哪些是语素基本音形，语素本身不能说明，还需要参考变音规则。我们给出一个确定语素基本音形的变音原则：

在语素组合中，相同的语素基本音形相加必然有相同的变

音形式。

根据这一原则,上声的几个可观察形式中,214 应该是语素基本音形。在不违背语素变音原则的前提下,一个语素独立分布的音形可以作为基本音形。如果和语素变音规则冲突,必须坚持变音规则,因为自由形式有时候并不一定是基本形式。比如大家经常提到的俄语例子:

	阳性、单数、主格	阳性、单数、所有格
葱	luk	luka
牧场	luk	luga

尽管 lug 是黏着形式,这时应该把 lug 作为牧场的基本形式,才能反映“葱”和“牧场”的对立。汉语方言中也有类似的例子。

如果一个变音行为对相同语素音形是周遍的,则该变音行为是规则变音,基本语素音形可以通过变音原则得到判定。比如上声变调。“一”的变调对 yi¹ 这个音形不周遍,所以词汇库中要有两个变体。

以上分析进一步说明,基本语素音形的确定和规则是互为条件的,基本语素音形不可能独立于变音规则而确定。

前面说,凡是规则变音都不必作为语素音形。Chomsky 和 Halle(1968,p11—12)讨论过不可预测的语音信息才存放在词库中。比如:

cat	cats
I	we

第一行只需要 cat 在词库中就行了。第二行的两个形式都必须在词库中。又比如:

telegraph	telegraphy	telegraphic
-----------	------------	-------------

只要词库中有 telegraph 的读音,其他两个词的读音都可以推导出来。这些结论都相当有道理。

前面讨论过,上声的 35 音值可以用规则推出,这种推导是以 214 为核心单位。当然我们也可能设想另一种可能,把上声的 35 音值作为基础,214 是推导出来的。但是,正如前面已经分析过的,这样做会把上声和阳平混在一起。这里已经引出了一个理论上的根本问题,语素音形的对立是不可能绕开的,否则就得不到基本语音单位或底层语音形式。从这种意义上说,包括生成音系学在内的所有音系理论都不可能绕开对立。有对立的音必须处理成不同的核心单位,否则就会出现相同的音(35)生成不同的音值的现象。

前面还提到,在大部分情况下,我们都可以用某个可观察形式作语素的基本形式,并且认为能够和其他语素区别开的形式是基本形式。助词“了”没有独立的声调,这时候需要抽象形式,由于轻音不仅仅在“了”中出现,所以需要统一起来,“了”表示为 le,不带声调,其他轻音语素也不带声调。当然也可以用其中的一种变调形式作基础,但不能反映轻音的本质。轻音的本质是自己没有独立的声调。Chomsky(1964)提出过抽象形式的问题。Chomsky 和 Halle(1968)以及后来的发展仍然肯定抽象的底层单位的必要性。这种思想和结构语言学的实证思想是完全不相容的。不过完全不考虑抽象形式有时候会遇到困难,考虑下面情况:

geology	geologic
telegraphy	telegraphic

如果我们以左边的实际读音为深层形式,下划线的部分由于处在弱化音节,都读[ə],通过规则生成右边形式时,这个[ə]变成了不同的读音,一个读[ɔ],一个读[æ],这时就没有条件可以解

释这种变化。另一方面,如果以右边的形式作深层形式,同样由于处在弱化音节,下划线的部分读[ə],通过规则生成左边的形式时,一个读[ɔ],一个读[e],条件也得不到解释。因此,以上词形必须有一个不依赖于实际读音的抽象的形式,其他弱化形式和表层形式都可以通过这个抽象形式根据语境派生或推断出来。这正是抽象形式得以成立的重要理由。

这种情况并不是孤例。考虑 Clark 和 Yallop (1995, p403) 的一个实例:

melody	melodic
heretic	heretical
demon	demonic
telephone	telephonist

以上两组形式都有弱化的央元音,只是部位不同。由于弱化的央元音是其他元音变来的,用以上两种表层形式作基本形式都会遇到困难,需要承认一个抽象的形式,以上两组形式都是抽象形式派生出来的。这是生成语法 underlying form 的一个重要证据。其实以上实例类型在很多语言中都存在。

不过,抽象的底层形式和可观察事实并不是没有关系,在大多数情况下,如果我们把语素基本音形作为音系分析的起点,抽象形式和可观察形式的关系可以得到更好的解释。前面说过,语素音形的区别性表现形式是基本形式,以上的英语的词例都分别由几个语素组成,这些语素都不单说,但有重读和弱读的分别,弱读的时候元音没有差别,重读的时候元音出现差别,重读时有差别的元音形式就应该看成基本形式。因此在深层形式中,每个语素都应该使用重读形式。最后抽象的词音形不过是每个语素重读音形在具体语境中的体现。由此可见,词音形中的抽象形式,从语素音形看仍然有可观察形式的基础,那就是语素音形的区别

性变体,严格地说就是语素基本形式。由于以上的每个词有两个语素可以有重读和弱读形式(当然每个语素的弱读都有可以预测的环境),当其中一个语素的某个音节重读时,另一个语素的某个音节处在弱读的状态,所以整个词才需要抽象的不可观察形式。以 telegraphy 和 telegraphic 为例,当 le 重读时,graph 弱读,当 graph 重读时,le 弱读,所以 telegraphy 的具体音值不能作为抽象形式。英语的字母拼写在多数情况下正好满足了把语素基本形式区别开这一点。所以在多数情况下,正像 Chomsky 和 Halle 所说,英语的字母确实可以作为底层形式。

确定语素音形以后,词或语符音形并不是完全可以通过语素音形预测,在英语中很明显。汉语中大部分词音形可以通过语素音形和规则预测,但轻声的分布不完全能预测。比如:

重轻	中重
汉语	汉族
蚊子	蚊帐
茶叶	茶杯

这些含有轻声的词需要词典标注。因此在词库中,还要有词音形的信息。在此基础上,有的组合的轻重现象可以预测:

X+重轻→中重轻:黑芝麻、古汉语、
重轻+X→中轻重:芝麻糊、文字屎

哪些能够预测,需要深入研究。不过语素音形是最核心的一层,和词音形不同,因为语素组合无论是词还是短语,其音形和语素内部音形的分布都有根本的区别。

通过以上分析可以看出,要充分提取核心语音单位,必须要对语素音形有全面的记录,列出同音语素音形表,而不仅仅是音

节表,也不仅仅是单音节语素表,因为语素音形可以是多音节的。我们发现,现在有好些音系描写只给出音位、音位变体、音节表和音位的组合规则,没有给出完整的语素音形表,这样的描写还不具有充分性。有的甚至没有音节表,更不具有充分性。语素音形表具有很高的—致性,只要得到语素音形表,就可以提取各层面的核心语音单位。但如果只有音位、音位变体、音节表和音位的组合规则,不一定能够复原语素音形,因此不一定能够充分提取核心语音单位。

同音语素音形表必须在田野调查中确定,因为语素音形的同一性需要和发音人当场核对。从这种意义上说,在田野调查中只记录大量词汇、句子和故事,然后离开发音人来分析语素和音位的方法会遇到无法进行同音语素音形核实的问题。当然,提取语素音形只是音系分析的起点,并不能代替音系、语流音变、词汇、语法的研究,所以记录大量词汇、句子和故事也是田野工作必不可少的。

把语素音形作为提取核心语音单位的起点,从某种程度上证明了汉语方言调查从单字字表出发基本上是合理的,因为在汉语中,一个字基本上是一个语素,像“盖儿”这样的儿化词都写成两个字,所以从单字字音入手的办法基本上等于王理嘉(1988)从单音节语素音形入手的办法。可以说,汉语从单字同音字表入手提取核心语音单位是一种比较有效的方法,因为字为语素音形提供了原形(陈保亚,1999)。王洪君(1999,2.2.5)认为汉语音系分析必须先分析单字音,或者说在不考虑叹词和语气词的情况下,这类单位构成的音系格局相当整齐(王洪君,1999,2.2.3),在演变方式上也有自己的内在规律(王洪君,1999,2.2.4),我们前面的分析基本上支持这样的观点。当然,有些语素在方言中是无字的,在汉藏语言调查中,很多语言也没有文字,另外,有不少语素不是单音节的,所以要充分提取核心语音单位,最终必须穷尽语素音形。也就是说,在

坚持单位和规则相互制约的前提下,语素音形一般是单字音,但单字音不一定总是语素音形,所以音系分析从根本上说要从语素音形入手。

既然语素音形是提取核心语音单位的起点,那么在田野调查中,从话语片断中提取语素就是首先要完成的工作。这对汉语和汉藏诸语言是比较容易的,但世界上有些语言,比如屈折变化很复杂的语言,有元音和谐的语言,语素的提取有一定的难度,在这样的语言中提取核心语音单位是否必须先从语素音形开始,需要做进一步的研究。

6.9 平行周遍条件与语音规则

先让我们先来考虑汉语上声变调的本质:

买鸡	买羊	买狗	买鹿	买
mai ²¹ tɕi ⁵⁵	mai ²¹ iaŋ ³⁵	mai ³⁵ kou ²¹⁴	mai ²¹ lu ⁵¹	mai ²¹⁴

宰鸡	宰羊	宰狗	宰鹿	宰
tsai ²¹ tɕi ⁵⁵	tsai ²¹ iaŋ ³⁵	tsai ³⁵ kou ²¹⁴	tsai ²¹ lu ⁵¹	tsai ²¹⁴

烤鸡	烤羊	烤狗	烤鹿	烤
k'au ²¹ tɕi ⁵⁵	k'au ²¹ iaŋ ³⁵	k'au ³⁵ kou ²¹⁴	k'au ²¹ lu ⁵¹	k'au ²¹⁴

“买鸡”的“买”([mai²¹])和“买狗”的“买”([mai³⁵])语音音质相同,声调有差异,我们不必把[mai²¹]和[mai³⁵]归并成语素变体,因为“买”和“宰、烤”形成一种平行变音,而且这种变音对每一个上声音节语素来说都是周遍的,可以概括成语音公式:

$$214 + 214 \rightarrow 35 - 214$$

$$214 + X \rightarrow 21 - X \quad (X \text{ 不等于 } 214)$$

$$214 \rightarrow 214 \quad (\text{单说})$$

再来观察另一类情况：

七张	七台	七两	七个	七
tɕ'i ⁵⁵ tɕaŋ ⁵⁵	tɕ'i ⁵⁵ t'ai ³⁵	tɕ'i ⁵⁵ liaŋ ²¹⁴	tɕ'i ⁵⁵ kə ⁵¹ tɕ' ³⁵ kə ⁵¹	tɕ'i ⁵⁵

八张	八台	八两	八个	八
pa ⁵⁵ tɕaŋ ⁵⁵	pa ⁵⁵ t'ai ³⁵	pa ⁵⁵ liaŋ ²¹⁴	pa ⁵⁵ kə ⁵¹ pa ³⁵ kə ⁵¹	pa ⁵⁵

“七、八”的语素音形变化也是平行的，可以概括成公式：

$$55+51 \rightarrow 35-51 / 55+51$$

$$55+x \rightarrow 55-x \quad (x \text{ 不是 } 51)$$

但这种平行行为只适用“七、八”两个特例。与此不同，前面讨论的上声语素“买、宰、烤”等的变化，可以周遍到每个上声语素。我们说上声语素“买、宰、烤”的变调行为在纯语音层面是平行并且周遍的，即只要一个语素的读音属于上声音节，就会变调，而“七、八”的变调行为在语音层面只平行但不是周遍，即不是所有的阴平调都有这样的变调。根据这一点可以断定“买、宰、烤”等的变调规则和具体的意义没有关系，而“七、八”的变调规则是和具体意义相关的，其他单念时读平声的单音节语素并没有这种变调行为。

再考虑下面实例：

阴平前	阳平前	上声前	去声前	单用	其他语子后
一张	一瓶	一两	一个	一	第一

i ⁵¹ tɕaŋ ⁵⁵	i ⁵¹ p'in ³⁵	i ⁵¹ lian ²¹⁴	i ³⁵ ke ⁵¹	i ⁵⁵	ti ⁵¹ i ⁵⁵
不歪	不斜	不好	不大	不	
pu ⁵¹ wai ⁵⁵	pu ⁵¹ ɕie ³⁵	pu ⁵¹ xau ²¹⁴	pu ³⁵ ta ⁵¹	pu ⁵¹	

“一、不”变调行为不完全平行,表现在单用的时候不一样,更谈不上周遍。这种变调方式也是和具体意义相关的。如果考虑到语法因素,“一”的变调更复杂(宋作艳,2005,1)。

以上三种变音行为可以概括成三种情况:

1. 独立变调:“亚”属于这种情况,变调行为是独立的,找不到平行的实例。

2. 平行变调:“七、八”代表了这种情况,两者的变调有平行关系;“一”、“不”作为两字组的前字,变调行为也有平行性,但单说时声调不同。

3. 平行且周遍变调:“买、宰、烤”所反映的汉语单音节上声语素属于这种情况。

第1种情况是没有规则的,其变调方式要在语符库中注明,这是变体处理的方式。第2种情况是准规则的,由于“七、八”这样的平行变调实例比较少,一般应该在语符库中解释变调行为,这也是变体处理方式。如果这种平行但不周遍的变调行为实例比较多,也可以放在规则库中解释变调行为,但语符库中仍然要注明哪些语素具有这种变调行为。第3种情况是平行周遍的,只需要在规则库中注明变调规则,不必在语符库中作任何解释和标注。以“买”为例,我们不说“买”有三个变体,因此在语符库中不必标注三个变体读音,只需要记录[mai²¹⁴]这个读音。所有的读音都可以通过特定语音条件用规则确定。这里的语音规则和具体意义无关。

汉语普通话中有一类后黏着的轻音语素,也是有规则的。比如

“子、头、了、着”等轻音语素,在不同声调的音节后面有不同的音高:

阴平 + 轻音	阳平 + 轻音	上声 + 轻音	去声 + 轻音
pau ⁵⁵ tsɿ ² 包子	p'iŋ ³⁵ tsɿ ³ 瓶子	fu ²¹⁴ tsɿ ⁴ 斧子 ^①	k'u ⁵¹ tsɿ ¹ 裤子
ɕin ⁵⁵ t'ou ² 心头	lai ³⁵ t'ou ³ 来头	ku ²¹⁴ t'ou ⁴ 骨头	y ⁵¹ t'ou ¹ 芋头
k'u ⁵⁵ lə ² 哭了	lai ³⁵ lə ³ 来了	tsou ²¹⁴ lə ⁴ 走了	tsuo ⁵¹ lə ¹ 做了
pau ⁵⁵ tɕə ² 包着	na ³⁵ tɕə ³ 拿着	lou ²¹⁴ tɕə ⁴ 搂着	faŋ ⁵¹ tɕə ¹ 放着

这类语素的音高变化也是平行周遍的,也是可以用语音规则预测的:

$X^{55} Y \cdot \rightarrow X^{55} Y^2$ (X:不包括声调的音节,Y:轻音节)

$X^{35} Y \cdot \rightarrow X^{35} Y^3$

$X^{214} Y \cdot \rightarrow X^{21} Y^4$

$X^{51} Y \cdot \rightarrow X^{55} Y^1$

这类后置的轻音语素元音有弱化央化趋势,都呈现出平行周遍性,因此,对这类有规则现象,语符库中只要存放轻音节的信息,不必存放每个轻音节的具体变体,然后在规则库中存放轻音变化规则。

汉语中还有一类轻音,和上面无自主声调的轻音要重一点。比较:

包子	瓶子	斧子	裤子
抓了	拿了	取了	送了
伸出	拿出	取出	运出
抓来	拿来	取来	送来
山里	楼里	水里	地里
新手	国手	水手	臭手

① 这时的上声“子”实际读音是半上。

续表

说者	仁者	老者	记者
山上	楼上	水上	地上
山下	楼下	水下	地下
爬下	拿下	取下	跪下

第一、二行的第二个字是前面提到的没有自己本调的轻音字，其他几行的第二个字轻音方式不同，和原调有一定的关系，但也都是有规则的，不必分析成变体，因此也不必存放在语符库中。

再考虑一个重庆话的实例。重庆话两个名词性语素重叠并儿化后，后一个音节的声调要变化：

55—55	31—55	51—51(31)	11—55
星星儿	头头儿	眼眼儿	棍棍儿
圈圈儿	瓶瓶儿	米米儿	盖盖儿
包包儿	帽帽儿	口口儿	洞洞儿
兜兜儿	皮皮儿	拐拐儿	柜柜儿
猫猫儿	毛毛儿	粉粉儿	豆豆儿
灯灯儿	铃铃儿	狗狗儿	岫岫儿
根根儿	雀雀儿		

这里的变调也是平行周遍的，凡是名物性可儿化单音节语素，都有这种变调行为，所以这种变调行为也是规则的，只是这种规则不是纯语音的，已经包含了“名物性”、“儿化”的语法信息。如果语符库中已经标注了语符的名物性特征，只需要在语音规则库中给出变音规则就行了。如果都用语素变体来描写，是相当复杂的。

如果坚持结构语言学的方法来提取语素，不明确区分规则和不规则问题，用变体来处理语素音形的变化，会使描写相当复

杂,看不出语素变音的活动规律。再以英语复数语素为例。英语复数语素有以下分布特征:

s 变体	出现条件	实例
[iz]	有擦音特征的辅音后	glasses
		roses
		dishes
		garages
		churches
		bridges
[z]	不具有擦音特征的浊音后	saws
		boys
		ribs
		sleeves
		pens
		hills
		cars
[s]	不具有擦音特征的清音后	books
		cliffs

按照 Bloomfield(1933, 13. 4, p211)和 Harris(1942)的分析,排除一些例外情况,英语名词复数可以作以下描写:

glass: glasses [iz]
pen: pens [-z]
book: books [-s]

即名词复数语素 s 有三个语素变体。前面讨论过,按照这种变体分析方法,汉语“买”的三种读音就会被分析成三个语素变体(morpheme alternants)。这种分析方法的弱点是,每个单音节

的上声语素都有三个变体。而按照音变规则处理办法,每个单音节的上声语素都只有一个独立的形式,在语音层面有两条规则,确定上声音节在语流中的读音。

其实 Bloomfield(1933, 13. 4, p211)已经看出了复数 s 的几种交替形式有下面的规律:

1. the alternation is regular (交替是规则的)
2. the alternation is automatic (交替是纯语音层面的)

但是 Bloomfield 仍然把 [iz]、[z]、[s] 看成是复数语素 s 的三个语素交替形式,并没有充分利用变音规则来处理变体。

从更高的层面看,把语素音形的规则变化处理成语素变体还存在更多的问题。我们联系英语的复数语素 s 变音和其他变音来进一步说明这个问题。和英语复数语素音形规则变化相关的还有动词 is 的减缩形式、动词第 3 人称单数 s、所有格 s。从语素交替的角度看, is 的减缩形式 's 有三个语素交替形式:

Bess's ready [iz]

John's ready [z]

Dick's ready [s]

第 3 人称动词的 s 也有三个语素交替形式:

He losses [iz]

He comes [z]

He works [s]

所有格 s 也有 3 个交替形式:

Marx's works ['mɑ:ksɪz 'wə:ks] [ɪz]
 children's books [z]
 my wife's sister [s]

如果把这些不同种类的情况都处理成语素交替形式,必然使语素变体大量增加。但是我们注意到,以上出现在其他 X 成分后的 s,无论是哪种情况,读音都只受语音层面规则的控制,并且是平行周遍的。可以给出这样一套规则:

X+s 的纯语音规则:

在擦音的情况下:ɪz

否则,在浊音的情况下:z

在清音的情况下:s

这种情况和前面分析的汉语上声声调交替情况相似,用分布会使语素变体增加,用规则可以简化描写。

这里的关键问题仍然是什么叫语音规则。一般地说,如果语素组合产生变音能够找到平行周遍条件,都可以作为规则来处理。前面分析过 Harris (1942, p109) Tübatulabal 的实例,其实也是规则描写的一个实例。

6.10 可延展性与音节的判定

音节是语言音响链条中的自然单位,是母语者最容易感知到的两种语音单位之一,另一种是韵(rhyme)。但是,后结构语言学和生成音系学早期的 SPE 理论基本上抛弃了音节这一单位。在后生成音系学理论中,音节的地位又得到恢复。音节这一单位确实是必要的,否则不能说明很多语言规则。比如:

VNN

碎纸机 * 粉碎纸张机

修车厂 * 修理汽车厂

NVN

纸张粉碎机

·汽车修理厂

这里语序转换的机制是什么,引起了广泛的讨论,但规则是明确的,即音节的数量在起作用,“碎、纸”都是单音节,采取 VNN 形式,“粉碎、纸张”都是双音节,采取 NVN 形式。又比如:

老陈

老张

老李

? 老端木

? 老欧阳

哪些姓氏可以进入“老 X”,音节的数量也在起作用。这说明很多语言规则都依赖于我们对音节的判定。

上声变调也是以音节为条件的,无论是单语素音节还是多语素音节:

手小

碗儿小

“碗儿”尽管是两个语素,但仍然是一个音节,上声变调在这个音节上发生,而不是在两个语素上发生。音节概念在后来的非线性音系学中有很重要的作用,很多音系规则都以音节的识别和音节的性质为前提。

但是,判定音节是一个很复杂的问题,前人已经提出了一些判定原则,解决了一些问题,但并没有从根本上解决问题,这是后结构语言学和生成音系学 SPE 理论抛弃音节的主要原因之一。由于音节在解释很多音系规则时有很重要的作用,非线性音系学从更高的理论角度重新启用了音节,并且给音节以非常重要的地位,但仍然没有给出判定音节的标准,只得承认音节是一个非实体单位。既然非线性音系学有很多音系规则依赖音节

概念,又不能给音节一个严格的定义或判定标准,这使得很多工作不能深入下去。音节并不是非实体单位,其实体性比音段要强得多。在语言调查中我们注意到,母语者判定一个音段序列有多少音节是比较一致的,尤其是在做诗或编歌词的时候能够断定两个音段序列的音节数是否相同,在给音符填词的时候能判定音节的数量(关于这一点后面还要讨论)。

Saussure 很早就注意到了音节的重要性和比较容易辨认的性质。Saussure(1916)说:

这种音位学的方法有一点是特别欠缺的:它过分忘记了在语言中不仅有一个个音,而且有整片说出的音;它对于音的相互关系还没有给以足够的注意。我们首先接触的不是前者,音节比构成音节的音更为直接。我们已经看到,有些原始的文字是标记音节单位的,到后来才有字母的体系。(Saussure, 1916, p81)

从田野调查的情况看,一个语言的音节数量是相当明确的,但音位、音段、音子则由于记音宽严不同或归纳音位的宽严不同,数量很不确定。这一点在前面已经有讨论。

音节的概念如此重要,又有如此一致的认同基础,有必要给音节一个严格的判定标准或定义。我们认为判定音节的困难主要有四个方面:

1. 一个音段序列中有几个音节?
2. 每个音节的核心是什么?
3. 音节中音段的分布顺序是什么?
4. 音段序列中的分界线在哪里?

以上 4 个问题是相互关联的,但不是同一个问题。如果断定一

个序列有几个音核,就等于断定了有几个音节。反过来,如果断定了有一个序列有几个音节,也就断定了有几个音核。不过断定一个序列是一个音节并不等于锁定了该音节中音核的位置。这需要解决第3个问题。前面3个问题即使得到解决,也并不等于第4个问题得到了解决。下面先讨论已有标准在解决上面4个方面的问题时的得失,再提出我们的一些解决办法。

一个音节通常都包含有元音,所以古希腊和古印度的语法学家提出了音节元音说的定义,大致认为一个元音代表一个音节,但是这个定义是模糊的,因为两个、三个元音的直接连接也可以是一个音节,也可以是两个或三个音节。比如两个元音:

语言实例	元音数量 2	音节数
爱	ai	ai 1
阿姨	ai	a,i 2

汉语中通常能够断定两个元音是一个音节还是两个音节,可以用声调来判定,但这不是一个原则。下面的实例[ɛian](西安)就不是靠声调可以断定的:

语言实例	元音数量 2	音节数
先	ɛian	ai 1
西安	ɛian	ɛi,an 2

在没有声调的语言中,判定音节更不能利用声调。如果像 Bloomfield 那样把[ai](爱)这样的音节看成只包含一个元音,或许可以为古希腊和印度语法学家的说法提供一种支持,不过这就走向了什么叫“一个元音”的术语之争。另外流音 l、r 和鼻音 m、n、ŋ 等非元音也可以成音节,这是元音说没有办法解释的。

Stetson(1924)提出过音节搏动理论,认为每个音节对应于一次胸肌搏动(a pulse of chest muscle activity)。Ladefoged(1967)对胸肌活动的调查并没有证明每个音节必然有一次胸肌

搏动。我认为大部分辅音元音交替连接的序列中,音节勉强能通过这个办法测试出来,比如英语的 worker,汉语的 gongren (工人),但对于 doing、lier 这样有两个或两个以上元音连接的片断,就难以判定,因为这样的元音连接的搏动方式和 lion、hair 的搏动方式很相似的。doing、lier 公认是两个音节,lion、hair 公认是一个音节。这里的关键问题是可成音节的音段在一起的时候怎么确定音节数量。再比如:

meal middle

这两个序列似乎都有两次胸肌搏动,但一般认为 meal 是一个音节,middle 是两个音节。如果坚持说 meal 只有一次搏动,就需要严格定义胸肌搏动的定义和制定严格的观察数据标准。汉语普通话的“裤子”算几次搏动?算几个音节?也存在类似的问题。其实胸肌搏动这个概念本身就是不明确的。

和搏动理论相似的一种理论是更早的时候 M. Grammont (1895)从发音角度提出发音器官肌肉紧张理论,发音器官的肌肉一次紧张和一次松弛形成一个音节,紧张点叫音峰,是音节的中心,通常是元音,松弛点叫音谷,是音节分界点。该理论的支持者苏联学者 Шерба, П. В. (谢尔巴)于 20 世纪 50 年代将该理论介绍到中国,对中国语言学界确定音节有比较大的影响,中国很多教科书都在一定程度上接受了这种理论。紧张理论本质上和搏动理论相似,只不过胸肌搏动是以发音器官肌肉的松紧交替的面貌出现的。肌肉紧张理论遇到的问题和搏动理论遇到的问题相似,不能判定几个元音音段的连接有几个松紧交替弧形,因此不能判定有几个音核、几个音节。和搏动理论一样,紧张理论也不容易判定快速读音中一个音段序列有多少紧张点,通常要求在慢速中确定紧张点。但汉语“有”这样的音节,在慢速中可能测试出 1 个紧张点还是 2 个紧张点?紧张理论也不容易判

定音核的部位,比如重庆话的 tɕiu“局”,不好断定是 i 紧张还是 u 是紧张点。另外松紧理论和实际观察到的事实也有出入,松紧理论认为音峰在元音上,比如 ta 中的 a 是音峰,因为 a 是紧张点,但实际上紧张点应该是在 t 破裂之前。

与松紧理论等价的是音响能量峰,即认为能量峰是音核所在。就像紧张点所遇到的困难一样,测定到多少能量峰往往是有困难的。

Saussure 认为内破和外破是音节区分的基本条件(1916, p54)。Saussure 的内破和外破是指和下面相类似的情况:

appa

alla

auua

aiia

aeaa

以上每个序列第 2 个音段是内破,第 3 个音段是外破。Saussure(1916, p90)说:

在音链中,由一个内破过渡到一个外破(>|<)就可以得到标志着音节分界的特殊效果。

但是怎样定义内破和外破,并不是很清楚。Saussure 的内破和现在所说的内破(implosive)并不完全是一个概念。另外, Saussure 判定内破实际上已经以音节的预先判定为先决条件,所以 Saussure 的论证存在循环。如果今后在不利用音节信息的条件下提出一套判定内破和外破的原则, Saussure 的理论或许会有用处。

目前影响比较大的是响度理论。早在 19 世纪末,欧洲学者讨论过响音和音节的关系。Jespersen(1913, 13. 12, p191)

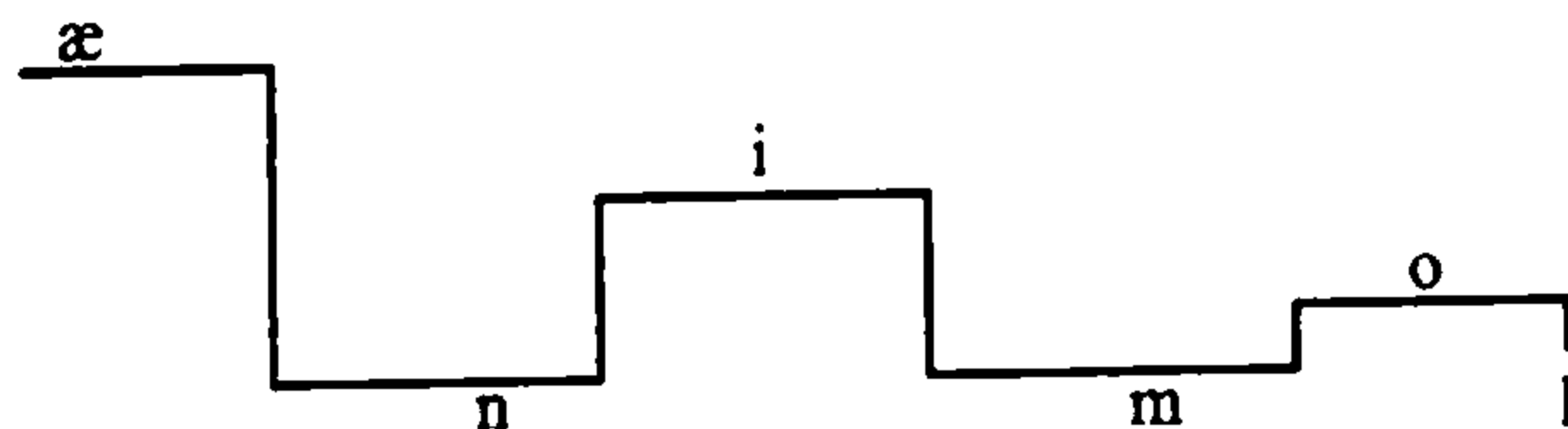
是根据响度判定音节的早期代表人物,他把语音的响度分成 8 个等级:

- | | | |
|-----|-----------|------------|
| (1) | 不带声(a)爆发音 | [p,t,k] |
| | (b)摩擦音 | [f,s,ʃ,x,] |
| (2) | 带声爆发音 | [b,d,g] |
| (3) | 带声摩擦音 | [v,z,ʒ,ʁ] |
| (4) | 带声 (a)鼻音 | [m,n,ŋ] |
| | (b)变音 | [l] |
| (5) | 带声各种 r 音 | |
| (6) | 带声闭元音 | [y,u,i] |
| (7) | 带声半开元音 | [ø,o,e] |
| (8) | 带声开元音 | [ɔ,æ,a,ɑ] |

其中数字越大,响度越大。Saussure 对这种理论提出过置疑 (Saussure, 1916, p83, p92)。Palmer, L. R. (帕默尔, 1936) 在 Jespersen (1913, p27) 的基础上给出了一个音节定义,是相当有解释力的:

代表响度的峰的就是音节的音。

根据这一定义,Palmer 断定英语 animal 有三个音节,因为有三个响度峰:



Jespersen 和 Palmer 的工作基本上奠定了音节响度理论的基础。后来的学者把音节响度理论更细化了,尤其是区别了核

心音节和边际音节,解决了很多早期响度理论遇到的例外。这方面讨论最深入的是响度顺序原则(sonority sequencing principle)。响度顺序原则最早由 Selkirk(1984)加以表述,他认为音节的核心响度最高,往两端的响度依次降低。响度顺序原则不容易量化,一般公认为响度顺序由高到低为元音、半元音、流音、鼻音、阻塞音。Clements (1990)提出了更细的标准,但问题还没有最终解决。不过元音作为核心,应该是没有问题的。

概括地说,音节响度理论认为音段的响度是不一样的,是有响度阶(sonority scale 或 sonority hierarchy)的,每一个响度峰(sonority peak)决定一个音节,每一个音节中有一个音段构成响度峰,前后是响度峰降低的音段,离响度峰越远的音段,响度越低,形成一个响度顺序。

响度顺序原则在一定程度上解决了一个音节内部音段的分布规律,但在判定一个音段序列有几个音节的时候,尤其是几个元音连接的序列有多少音节,这种理论遇到的问题和搏动理论、紧张理论遇到的问题相似。比如[ai],可以有两个音核,也可以只有一个音核,这是响度原则无法解决的。即使能够判定响度峰,有时候也不能判定有一个音节。比如英语 excuse [ɪkskjʊ:z],其中[ksk]形成一个响度弧形,s是一个响度峰,应该有一个音节,但英语母语者通常认为 excuse 只有两个音节。有可能母语者判定音节数量的标准不是建立在响度基础上的,因为从一串耳语发音中,我们仍然能够判定音节数量。比如,说出一句耳语“玩儿扑克牌儿”,我们仍然能够断定有4个音节而不6个音节。由于说的是耳语,肯定不是依据响音理论来判定音节数量。另外,即使在解决音节内部音段分布规律上,响度原则也还存在问题。响度的量化目前还停留在听觉上,声学参数还没有确定下来,因此不同的学者响度阶并不完全一样:

响度阶	Jespersen	Durand
1	不带声(a) 爆发音 [p,t,k]	清爆破音
2	(b) 摩擦音 [f,s,ʃ,x ,]	浊爆破音
3	带生爆发音 [b,d,g]	清摩擦音
4	带生摩擦音 [v,z,ʒ,ʁ]	浊摩擦音
5	带声 (a) 鼻音 [m,n,ŋ]	鼻音
6	(b) 变音 [l]	l 类音
7	带声各种 r 音	r 类音
8	带声闭元音 [y,u,i]	i,u
9	带声半开元音 [ø,o,e]	e,o
10	带声开元音 [ɔ,æ,a,ɑ]	a

以上我们是按照数字越大,响度越高的方式来排列 Jespersen 和 Durand 给出的响度阶,可以看出,斜体两行的响度正好相反。正因为如此,音段的响度分阶存在循环论证,即研究者有可能先根据音节中音段的分布顺序来确定响度的阶,然后又说响度的阶决定了音节中音段的分布顺序。响度原则另一个循环论证表现在对核心音节和边际音节的区分,响度顺序原则把符合响度顺序原则的音节看作核心音节,把不符合响度顺序原则的音节看作边际音节,于是说一个音节符合响度顺序原则是因为这个音节是核心音节,反过来又说一个音节是核心音节是因为符合响度顺序原则,并没有提出一个独立于响度顺序原则的手段来区分核心音节和边际音节。音节内部音段分布的阶或顺序可能是存在的,但似乎不是响度决定的,可能有另外的原因决定着音节内部音段的分布规律。顺序原则断定一个音节内部的音

段分布有一定效力,但响度顺序原则在理论上还不严格。

从以上分析可以看出,断定多元音序列有几个音核、音节,跟断定一个音节中的音段分布顺序,尽管有关系,但不是一个问题。目前最困难并且最迫切需要解决的仍然是音核和音节数量的判定问题。以上搏动理论、紧张理论、能量理论、内破外破理论、响度理论等都有一个共同点,或者是从肌肉活动的强弱出发,或者从响度强弱出发,统一起来都是从强弱上考虑音节,即认为音节是音段序列中强弱交替的结果,一个强弱交替代表一个音节。我们把这类观念统一称为音强音节观念。音强音节观存在两个根本的问题,一是不能充分的判定音核的数量,一是不能对耳语的音节切分做出解释。

正是因为母语者在耳语中也能断定音节的数量,我们认为音节的规则不应该在强弱中去找,而应该绕过强弱。有人曾用元音开口度大小来定义元音的响度。开口度越大的音,响度越大。比如 Selkirk(1984a, p112)提到元音开口度越大,响度越高,从 Jespersen(1913)所列出的响度阶看,也大体有了元音开口度越大响度越大的观念。Saussure(1916, p85)认为开口度最大的 a 不会有外破和内破的区分,实际上也暗示开口度的作用。用元音开口度大小来量化元音响度的主要目的之一是为元音响度找出可观察的量化标准,以便更好的说明响度顺序理论,尤其是说明类似 ai 这样的多元音构成的音节中哪个元音是音核。但是,上面我们提到,在耳语音中也可以判定音节数量,因此响度可能也不是音节的根本属性,继续用各种手段来定义响音等级以达到判定音节的目的在认识思路上存在问题,用开口度的大小来定义响度再来确定音核是多余的。不过,用开口度的大小来定义响度等级却给我们一个启示,开口度本身和强弱是没有关系的,我们可以直接从开口度的大小来观察音核的分布,这样就可以绕开长期以来困扰音节判定标准的强弱的视角。在汉语普通话中,多元音音节中开口度大的元音总是音核。比如:

带多元音的韵母	最大开口度音	韵腹	音节数
ua	a	a	1
ia	a	a	1
uo	o	o	1
ie	e	e	1
ye	e	e	1
ai	a	a	1
uai	a	a	1
ei	e	e	1
uei	e	e	1
au	a	a	1
iau	a	a	1
ou	o	o	1
iou	o	o	1
ian	a	a	1
uan	a	a	1
yuan	a	a	1
uen	e	e	1
ianŋ	a	a	1
uanŋ	a	a	1
ionŋ	o	o	1
uenŋ	e	e	1

汉语普通话多元音音节中开口度最大的元音就是韵腹,也是音核。断定有几个音核也就决定了几个音节。我们可以把这个原则称为音核判定的开口度原则。一般地说,元音序列的开口度有几次高低起伏,就有几个音核,也就有几个音节。开口度原则绕开了对响度的依赖,可以比较合理地说明耳语的音节。从某种意义上说,把汉语普通话“韦”的韵母写成[ui]而不是[uei],似

乎不能反映音节开口度的信息。

每个带有元音的音节都有一个最大开口度元音。从大量语言的情况看,三个元音的音节中都是最大开口度的元音居中。因此,在可以比较元音开口度的情况下,可以认为最大开口度元音是音节的核心,较低开口度的元音分布在最大开口度元音的周围。中国传统音韵纯元音音节的韵头、韵腹、韵尾中,都是韵腹开口度最大。一般地说,在可以判定开口度大小的情况下,开口度大的是核心。根据这一原则,一个音段序列至少有多少音节就可以判定。比如,[aeo]至少是两个音节,因为e开口度小于左右的元音开口度,不可能是这三个音段的核心。[oua]也至少是两个音节,因为u的开口度小于左右的元音开口度,不可能是这三个音段的核心。不过,开口度原则只能判定[aco]这样的音段序列中至少有多少音核和音节,不能判定最多有多少音核和音节,因此仍然没有从根本上解决音核和音节的判定问题。

对于一个确定的音节来说,如果该音节包含两个或两个以上的开口度不同的元音,开口度原则可以判定音核,但是在开口度相当或舌位高低相当的情况,这个原则不能贯彻下去。仍然以重庆话[tɕiu]为例:

局:tɕiu

i和u的开口度是相当的。这时候仍然必须借助后面提出的延展原则。根据延展原则,“局”的音形中u是音节核心。

即使一个多元音音节内部元音的开口度不同,开口度原则也不能够完全判定音核的位置。比如,英语的[hɪə](hear),ə的开口度比i大,但音核是i(后面会详细讨论判定标准)。类似的例子:

fear	[fiə]
sure	[ʃuə]

傣语中有[t'au³⁵], 音核是在 u 上, 后面还会证明这一点。前面提到重庆话中有[tɕiu]这样的音节, [i]和[u]的开口度是相当的, 但音核是在[u]上。英语还有三元音的情况:

hire	[haɪə]
hour	[aʊər]
loyal	[lɔɪəl]

由于这些三元音序列的中间一个元音开口度最小, 所以这些三元音序列都有两次高低起伏, 但是很多学者认为这些三元音序列只是一个音节。这说明开口度原则不能断定多元音序列中有多少音核和音节, 同时还说明即使在一个音节中, 开口度原则也不能完全断定元音的分布顺序。

开口度原则遇到的另一个根本问题是不能处理辅音音节, 比如边音、鼻音等是不好判定开口度的。法语中的音节 pst, 其中 s 是音核, 开口度也没有办法解释, 这一点决定了开口度只是带元音音节的一种性质, 还不是所有音节的根本属性。

很少有人从音长的角度来讨论音节的判定问题。下面我先从田野调查的角度提出一个延展原则, 然后再分析其他原则的适用程度。在语言调查中我们发现, 发音人通常能够判定两个序列的音节是否相同, 而且在诵读中音节核心可以延长, 在音节和乐曲音符的配合中, 音符一般只配音核。根据所调查和接触到的语言, 我们总结出一个延展原则:

每个音节是有一个且只有一个音段的口形是可延展的。口形可延展的音段是这个音节的音核。有几个音核就有几个

音节。

所谓口形可延展,是指在话语中可以延长读音,包括说耳语。比如汉语普通话的“爱”/ai/,口形可延展的音段只有[a]而不是[i],即只有[a]可延长读音(包括耳语):

a——i

所以“爱”是一个音节,其中[a]是音节的音核。通常情况下,我们把口形可延展的音段称为可延展音段。“鸭”/ia/的可延展音段只有[a]:

ia——

所以“鸭”也是一个音节,其中 a 是音节的音核。“歪”/uai/的延展音段也只有[a]:

ua——i

所以“歪”也是一个音节,其中[a]是音节的音核。

口形延展允许特征的传递,比如[san]的延展并不仅仅是[a]的延展,而是[ā]的延展,所以我们说是口形的延展而不是音色的延展。

音节的可延展和音节的长短是不同的概念。只要是音节总是可以延展的,但开音节和闭音节的长短在不同的语言中有比较大的差别。不过,根据我们的调查,即使是那些音长很短的闭音节,在特定的语用条件(不是人为发音条件)下都是可以延展的,比如傣语的“十”([ɕip³⁵])。

延展原则中音段的延展是指一次延展。通常情况下相邻的

两个音节的音段都是不一样的,不一样的音段通常有不一样的口形,所以有几次延展是容易判定的。当两个音节是两个相同的音段,这时有完全相同的口形,比如在汉语的乘法口诀中有“一一得一”,前两个“一”的口形可以说是完全相同的,这时的延展是两次延展,中间是有停顿的。当然,如果把“一”的音值理解成[ji],“一一”就成了[jiji],舌位有改变,更不存在判定问题。

既然两次延展中间可以有一次停顿,为什么不直接用停顿来判定音节数?实际操作起来停顿的判定有一定的困难。比如在我们的初步测试中,有人认为 meal 可以有两次停顿。但这个片断只可延展一次。

延展原则中的可延长性质都是指表达的需要出现的延长,不包括语言学家或研究者从研究的角度对音段故意拖长的情况,比如 rice 中,研究者可以故意把[s]发的得很长,以显示擦音的性质。但在自然话语的表达功能中,这里的[s]是不延长的。

歌曲中音符的延长和歌词音节的延长最能体现延展原则。根据我们的初步调查,在歌曲中,音符和可延展音段有下面的关系:

1. 可延展音段的出现都伴随音符的出现;
2. 音符的延长都在可延展音段上;
3. 音符变化或旋律的起伏都在可延展音段上。

比如下面的民歌《在那遥远的地方》第一句:

me	so	la	sofa	me	so	la	—	—	sofa
在	那	遥	远的	地		方			
zai	na	yao	yuande	di		fang			

这里的读音用普通话汉语拼音,旋律用首调唱名法记音。其中

汉语拼音黑色斜体字都是可延展音段,其他非黑色斜体字音段,不可能延长和作旋律起伏。因此,也可以把可延展音段称为音符音段。根据我们对汉语歌曲和中国民族语言歌曲的初步调查,旋律和音段的匹配都符合这种规律。

有一种衡量音节长短的莫拉(mora)理论,简单地说,长元音构成的音节或重音节通常是两个莫拉,短元音构成的音节或轻音节是一个莫拉,不过 kip 中的 ki 是一个莫拉, p 是一个莫拉。经过 Selkirk(1978, 1984)、Hayes(1984)、Halliday(1985)、Nespor 和 Vogel(1986)、Zec(1988)、张洪明(1992)等学者研究,韵律单位的层级由大到小分为若干层,其中音步(foot)以下包括音步、音节、莫拉(mora)。但是,莫拉理论本身不能判定音节数量,而是在已知音节数量的情况下断定音节的长短。从音节的可延展角度看,无论是一个莫拉的音节还是两个莫拉的音节,都是可以延展的,并且都只有一个可延展音段。汉语中的虚字,如“了、着、过、的、地、得”,尽管音长很短、音强很弱,但其中的元音都是可延展的,所以仍然是一个音节。这些虚字在歌曲中都可以落在音符上。比如大家熟悉的歌曲《在那遥远的地方》:

在那遥远的(*fa*)地方,有位好姑娘,人们走过了(*mi*)
他的(*mire*)帐房都要回头留恋地(*do*)张望

这里的斜体字都是虚字,其中“他的帐房”中的“的”,还可以有旋律起伏,即从音符 *mi* 到 *re*。

英语中也有很多多元音结构,下面带双元音的音段组合通常都被当做单音节:

[aɪ]/[ai]: buy, eye, high, ice, why, try, fly, tie, die, ei-
ther, height, aisle,

- [ej]/[ei]: pay, bay, hay, day, say, name, rain, obey, eight, break
- [ij]/[i:]: bee, fee, bleed, key, me, sea, see, even, people, machine, field, receive
- [uw]/[u:]: do, should
- [ɔj]/[ɔi]: boy, coy, toy
- [aw]/[au]: cow, how, bough
- [ow]/[ou]: low, sow, owe, toe, tow, hoe, go, show, flow
- [juw]/[iou]: radio

以上音节的可延展音段都是黑斜体音段。据我们的初步调查，这些音节在英语歌曲中的可延展音段也都是斜体的音段。以上可延展音段前后的音段也可以看成是[j]或[w]。根据这种看法，以上都是单元音音节。英语中还有一些真正带多元音的组合，比如：

	单音节	双音节
aiər/aiə	hire, ire, fire, lyre	higher, liar
ɛər/ɛə	fair, fare, care, dare, share, there, their, air	
auər/auə	flour, lour, coward, tower	flower, cowered
ɔər/ɔə	door, more, pour	
iər/iə	here, fear, ear, dear, deer, weird, clear	
uər/uə	shure, poor, moor, tour	

大多数学者都认为第二栏的组合是单音节,但也有些人认为其中有的是双音节,比如上面所标注的。到底怎么办还没有统一的标准。从延展原则看,有两种情况。以[$\epsilon\text{ə}$]在英语歌中的延长情况为例:

Scarborough Fair :

6	6	3 3	3	7 1 7	6	6—
.	.			.	.	.
Are	you	going	to	Scarborough	Fair	(f ϵ ə -)

这里[$\epsilon\text{ə}$]中的[ə]在6音符上延长。再看另一乐句:

6 6	6	5	3 3	2	1	7 5
						.
remember	me	to	one	who	lives	there (ð ϵ ə -)

这里[$\epsilon\text{ə}$]中的[ə]在5的起始位置出现,并且延长。再看下面的情况:

The sound of silence :

5—3

People writing songs that voices never share (share: sh ϵ -ə)

这时的延长音符是在[ϵ]。由于[$\epsilon\text{ə}$]中的[ϵ]和[ə]都可以延展,带[$\epsilon\text{ə}$]的组合似乎和带 ai 一类的组合不同,可以看成两个音节。以上带[ə]的组合也有的只有一种延展方式。仍然以英语歌为例:

Everytime :

i make believe that you are here. (here: hi-ə)

it's the only way i see clear. (cli-ə)

即[iə]中的[i]是延展音,[i]应该看成音核,这和前面提到的开

口度原则是矛盾的,和一般公认的响度原则也是矛盾的。

以上都是带[ə]的多元音组合的情况,看来[Xə]是一个音节还是两个音节还需要进一步研究。

根据延展原则所确定的音核和根据开口度原则所确定音核是不完全一致的。前面提到英语 hear[hieə],延展音段是[i],最大开口度音段是[ə]。傣语[t'au³⁵](哪里),延展音段是[ɯ],最大开口度音段是[a]。如果我们依照开口度原则,就不能解释歌曲中的[t'au³⁵]只有 ɯ 为最大延展这事实。

在汉语这样的声调语言中,多元音音节的可延展音段和声调的走向大致是:

	正常模式	延展模式	错误延展模式
阴平	5—5	5—	5—
	ma-u	ma-u	mau-
阳平	3—5	3—5	3—5—
	ma-u	ma-u	ma-u-
上声	21—4	2—1—4—	2—1—4—
	me-i	me-i	me-i-
去声	5—1	5—1	5—1
	le-i	me-i	me-i-

延展模式说明韵尾所占的时间格很短,而延展音几乎可以无限延长,相比之下,延展音段几乎占用了声调的全部调型。正常模式下,给人的感觉是韵尾落在声调的后一个音高上,这似乎说明正常模式中韵尾落在声调最后音高上只是音节延长时间不显著形成的听觉。错误模式证明韵尾不可能延展。上声特别能说明问题,如果延展特别长,以上韵尾 i 既不可能在 i 上延长,也不可以在 4 上延长。

根据延展原则,汉语普通话多元音音节的音核都可以得到判定。一般人把汉语普通话的韵腹作为音核,根据延展原则,韵

腹确实都是可延展的,以零声母多元音音节为例:

零声母多元音韵母	可延展音段	韵腹	音节数
ua	a	a	1
ia	a	a	1
uo	o	o	1
ie	e	e	1
ye	e	e	1
ai	a	a	1
uai	a	a	1
ei	e	e	1
uei	e	e	1
au	a	a	1
iau	a	a	1
ou	o	o	1
iou	o	o	1
ian	a	a	1
uan	a	a	1
yuan	a	a	1
uen	e	e	1
iang	a	a	1
uang	a	a	1
iong	o	o	1
ueng	e	e	1

因此,以韵腹作为音核和根据延展原则判定音核是一致的。从非自主音系学的角度看,可延展音段在音值上会不同程度的受到前后音段的影响,但可延展音段的可延长性质基本上是可以确定的。

根据延展原则,元音连接序列中有几个音节的问题可以得

到解决,判定方式是有几个延长点就有几个音节。当然,延展原则是指正常情况下对音段序列有多少音核和音节的判定。在实际语流中,两个音节快读时听起来像一个音节,比较:

西安去了两次
先去了两次

快读使这两句话听起来很相似:

先 xian 西安 xian

但这并不能阻碍延展原则对这两个序列有多少音节的判定。在延展测试中,“先”只有一个可延展音段,即鼻化的 a,是一个音节。“西安”有两个可延续的音段 i 和鼻化的 a,是两个音节。上面的英语实例 doing 有两个延展音段,所以是两个音节。根据同样理由,meal 是一个音节,middle 是两个音节。这里的延展可以归结如下:

延展序列 1	延展序列 2
meal	
mi	dle
do	ing

延展原则是在延长音段口形中判定音核,并进一步断定音节。前面总结说,搏动理论、紧张理论、能量理论、内破外破理论、响度理论等都有一个共同点就是从强弱上考虑音节,即认为音节强弱交替的结果。我们把这类观念统一称为音强音节观念。从延展原则的角度看,音强音节观也是有道理的,因为以延展音段为音核构成的音节链条中,确实能够从频谱上观察到振幅强

弱的交替,并且这种交替的数量基本上和延展音段构成的音节的数量是一致的。但是音强理论仍然不能作为判定标准,因为不通过延展的观察,强弱是观察不出来的,或者说如果不延长读音,就不可能观察强弱。或许有人说搏动理论和松紧理论也可以在慢速读音中分辨出音节来,但实际上测试起来有困难。比如

meal

middle

要断定 meal 有一次搏动,或一次紧张,而 middle 有两次搏动或两次紧张,是有一定困难的,而用延展,就很容易把两者区别开。另外,强弱理论不容易观察出音峰(音节核心),更难观察到耳语的音峰。

延展原则要求能成音节的音一定是可以延展的,而可延展的音也有资格作音节。由于元音都是话语中可延展的音段,所以都有资格作音节。辅音中有一部分不可延展,有一部分可延展。[m],[n],[ŋ],[l]等流音之所以能成为音节,是因为可以延展。[v]也可作音节,比如成都话的 v(五),因为[v]也是可延展的。其实带有擦音特征的辅音,由于气流可以持续,都是可延展的,都有理由成为音节,这样就能解释为什么法语的[pst](喂)可以是音节。也可以认为有些可延展辅音成音节时已经带上了一个清化元音或不带音的元音,比如法语[pst]中的[s]可以认为带上了一个[ɪ]的口形,无论采取哪种看法,都构成了可延展的条件,因此也具备了成音节的条件。

可延展音只具备了作音节的条件,至于是否成音节,需要看在具体的语言中是否可以延展。英语的 next 中有音段序列[kst],但英语中的[kst]并不延续,所以不是音节。延展原则中所说的可延展性都是指具体语言系统中的可延展性。语音学上的可延展性只是语言系统中可延展性的必要条件。

next 这个例子说明,即使响度峰的存在也不能必然断定音

核的存在,因为 next[nɛkst]中有 ɛ 和 s 两个响度峰,但 next 只有一个音节。前面讨论过英语的 excuse,由于其中有一个序列是 ksk,根据响度理论,形成一个响度弧,s 是响度峰,因此整个 excuse 有三个响度峰,应该有三个音核,有三个音节,但根据延展原则,ksk 中的 s 没有可延展性,即 s 不可以延长。

以上延展原则的定义中没有说“延展浊音”或“延展响音”,而只说“延展”.因为有的语言中,延展音段的存在就可以形成音核和音节,并不需要响音,比如在法语中,pst(喂,唤人注意)中的 s 是可延续的,pst 是一个音节,s 是音核,但 s 不是浊音。又比如汉语普通话 ɕy(嘘),也可以看成是一个音节。其中 y 是非浊音,只是 y 的口型动作,ɕy 是一个音节,音核是声带不震动的 y。

延展原则或许能够解释响度分阶的矛盾。前面提到,Jesperson 和 Durand 的响度分阶在下面两组音段上是有矛盾的:

	Jesperson	非线性音系学	Durand
2	不带声摩擦音 [f,s,ʃ,x,]	[b,d,g]	浊爆破音
3	带生爆发音 [b,d,g]	[f,s,ʃ,x,]	清摩擦音

即 Jespersen 认为 [f,s,ʃ,x,] 的响度低于 [b,d,g] 的响度,但 Durand 的看法正好相反,[f,s,ʃ,x,] 的响度高于 [b,d,g] 的响度。从延展原则看,[f,s,ʃ,x,] 是可以延展的,可以作音节核心,如上面提到的法语 pst(喂),但 [b,d,g] 不能延展,作音节核心的例子也没有,因此在音段排序上,[f,s,ʃ,x,] 比 [b,d,g] 的核心程度更高。回头再来观察响度阶的排序,可以发现,音段基本上可以分成可延续音段和不可延续音段,作音核的一定是可以延续的音段。

在很多情况下,根据元音和辅音的分布往往能够判定音节的数量。音节的一般模式是:

CmVfCn

其中 C 为辅音, m、n 表示辅音数, 可以为零。V 为元音, f 为元音数, 一般不为零, 除非 C 中有流音。f 最大值为 3。f 最大值为 3 的理由将在后面给出。由于 f 的最大值为 3, 最小值是 1, 所以当两个或两个以上的元音连接时, 需要由可延展原则判定是一个音节还是两个或两个以上的音节。汉语由于有声调, 一般不存在音节的判定问题。

以上的模式同时暗示, 一个音节中的元音丛不可能被辅音隔开, 所以下面格式一定是两个音节:

VCV

但是通过元音的分布来判定音节的个数这一原则不能贯彻到底, 因为鼻音、流音都可能成为音节。所以从根本上说, 音节数的确定仍然需要根据延展原则。

6.11 两类语音组合规则

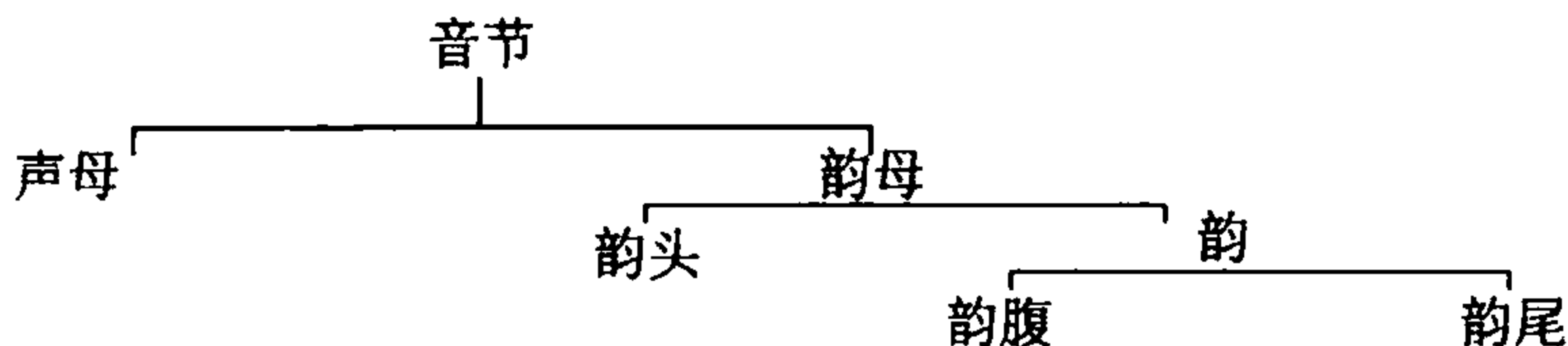
6.11.1 音段组合成音节

结构语言学注意到音位和音素的区别, 同时也注意到纯语音的音节和具体语言中音系的音节, 因此主张区分语音学的音节和音系学的音节。事实上我们不可能给出一个普遍的音节清单, 类似语音学家给出的国际音标清单, 从这种意义上说, 音节都是音系学上的, 每种语言的音节结构都不完全一样。另一方面, 判定音节和分析音节肯定要有语音学的参数, 因此每种语言的音节也都离不开有语音学的性质。要区分语音学的音节和音系学的音节往往涉及我们怎样严格定义语音学和音系学。

语音组合规则有两个重要的层面, 一个是音段组合成音节, 一个是音节组合成音步或更大的片断。

根据延展原则来判定音核,音节的数量也可以得到判定。现在的问题是,在一个音核的周围,音段的分布是否在全人类的语言中都有普遍的顺序。这是响度顺序原则要回答的问题。这也是我们前面提到的音节所面临的第三个问题。这个问题是否可解现在还没有得到证明。事实上很多实例显示音段分布并不一定有顺序限制,比如英语 box[baks](盒子)和 bosk[bask](小树林),都是核心音节,一个带有 sk 音段分布,一个带有 ks 音段分布。类似的还有 max(最大)和 mask(口罩)等。如果这类现象得不到解释,提出音段普遍顺序问题这一思路的合理性就会受到质疑。音段的分布顺序很可能只是语言系统的内部属性,比如汉语普通话中 p、t、k 不分布在音节末尾。下面我们先把音段顺序问题放在一边,重点讨论音段组合成音节的规则。

不同的语言或方言中,音段组合成音节的规则并不完全相同,有的甚至有很大差别,因此音段组合成音节的规则主要是语言或方言系统内部的规则。汉语普通话的音段组合规则已经得到了充分的表述。除了声调,汉语普通话音节一般可以分成声母、韵头、韵腹、韵尾:



也有人主张韵头和声母组成直接成分。由于每个音节都必须有一个音核,而汉语普通话的音核就是韵腹,所以一个音节的韵腹不能为零,其他部分都可以为零。也有人主张零声母前面实际上是有声母的,比如开口呼前面有一个 ? 声母,合口呼前面有一个 w 声母,齐齿呼前面有一个 j 声母,撮口呼前面有一个 y。根据这种认识,汉语普通话的声母也不能为零。

从每个非轻音音节的时长基本相同的角度看,也可以说

声母、韵头、韵尾都不为零,不同的音节只是选择音段的方式不同:

	声母	韵头	韵腹	韵尾	传统分析类型
啊	ʔ	a	a	a	零声母
无	w	u	u	u	零声母
衣	j	i	i	i	零声母
鱼	ɥ	y	y	y	零声母
家	tɕ	i	a	a	无韵尾音节
瓜	k	u	a	a	无韵尾音节
卡	k'	a	a	a	无韵尾音节
弹	t'	a	a	n	无韵头音节
天	t'	i	a	n	韵头韵尾齐全
团	t	u	a	n	韵头韵尾齐全
全	tɕ'	y	a	n	韵头韵尾齐全

这样的分析比较容易说明音节和音节的组合规律,比如前面讨论儿化时提到的内容。也有利于建立对应规则。当然,实际上这里音节的长短可能是不一样的。

音段组合成音节的规则是不周遍的,不周遍的程度有时候甚至要求一个音节给出一条规则,比如说“佛”([fo³⁵])。从这个角度看,音节内部的音段组合规则本质上是理解规则。

音段组合成音节的规则也有很大的相对性,这取决于我们怎么归纳音位。这一事实本身说明音位的问题不能绕过。

6.11.2 音节和音节的组合

如果承认音节的重要地位,除了切分音节,另一个重要问题是音系底层表达中是否需要音节。音节组构(syllabification)理论认为音系的底层表达不需要音节,音段是基本的,音节是通过音段和规则生成出来的(Goldsmith,1990)。这一问题从另一个

角度看就是是否需要区分基本音节和派生音节,音节组构理论实际上是不承认或不重视基本音节。我们认为应该区分基本音节和派生音节,有三方面的理由。首先,从语言分析程序角度看,音节应该是比音段更初始的概念,因为音段是在音节或基本语素音形的基础上通过对比分析提取出来的(见前面有关章节的论述)。其次,从音段生成音节的规则并不是周遍的,比如,汉语普通话 u 可以和舌面音以外的辅音组合,但这条规则并不周遍,仅以舌尖前声母为例:

租 tsu ⁵⁵	族 tsu ³⁵	组 tsu ²¹⁴	[tsu ⁵¹]
粗 ts'u ⁵⁵	徂 ts'u ³⁵	[ts'u ²¹⁴]	醋 ts'u ⁵¹
苏 su ⁵⁵	俗 su ³⁵	[su ²¹⁴]	速 su ⁵¹

方括号中的斜体音节只是潜在的,不是现成组合,词汇库中必须提供这方面的信息。区分基本音节和派生音节的另一个理由是派生音节是由基本音节生成的,不区分这两类音节就不能认识到新的音节的形成规则。基本音节是从基本语素音形中得到的音节。汉语普通话的[kar⁵¹](盖儿),是一个音节,但这个音节是由两个语素[kai⁵¹](盖)、[er³⁵](儿)组合形成的,因此不是基本音节,其中的[kai⁵¹]和[er³⁵]都是一个语素,[kai⁵¹]和[er³⁵]是基本音节。有了基本音节,所有的儿化音节都可以通过规则推导出来,前面已经做过分析。在基本音节的基础上通过规则推导出派生音节,可以使描写简化。一个语言的音节是可以在词汇库中给出的,但是如果不区分本音节和表层音节或变音节,单位库中的音节数量就会相当多。当然,这还不是承认基本音节的根本理由。承认基本音节的根本理由在于对规则的充分描写,因为儿化音节从根本上说不是通过记忆获得的。对于任何一个基本音节 X,北京人都能够知道 X 儿化后的读音,即使他从来没有使用或听说过这个儿化词。

前面提到,汉语的轻声音节实际上在不同的语素后面有不同的音高,比如:

飞了(lə³) 来了(lə³) 走了(lə¹) 去了(lə²)

并不需要把这些不同音高的[lə]作为不同的基本音节,只需要一个[lə]音节就够了,其他的音高都可以推导出来。这个规则不仅仅周遍[lə]音节,对于任何一个没有本调的轻音音节 X,这个规则都是有效的。这不仅使描写简单,更重要的是解释了一种规则。

因此,需要区分核心音节(或基本音节)与非核心音节(或派生音节)。核心音节是从基本语素音形中通过延展原则得到的音节。我们把英语 working 中的 ing 确定为一个核心音节,就在于先得到了 ing 这个语素,然后根据延展原则确定 ing 只是一个音节。有时候一个语素有两个可延展音段,比如汉语普通话[la⁵⁵ tɕi⁵⁵](垃圾),其中每个可延展音段都代表了一个音核,因此“垃圾”是两个核心音节。一般来说,一个语素有几个可延展音段就有几个核心音节。核心音节是另一个重要概念,一旦核心音节确定,其他更小的核心音系单位也容易确定。

在英语中,也需要区别基本音节与派生音节,否则无法认识语音规则。考虑下面的实例:

基本音节	基本音节		派生音节	派生音节
work	ing	→	wor(k)	king
sleep	er	→	slee(p)	per
list	ing	→	lis(t)	ting

我们可以概括出一条两个基本音节组合的规则:如果第二个音节是元音开头,第一个音节的最后一个辅音会移动到第二个音节的开头。英语中很多话语序列的表层音节都可以通过基本音节的组合规则推导出来。比如,英语中仅 ing 的组合就可能产

生大量派生音节:

ping←sleep+ing
 ving←give+ing
 king←work+ing
 ting←list+ing
 ding←read+ing

词音形和词组音形中的音节都应该是表层音节,很多是派生音节,基本音节都是在语素音形中确定的。

音节组构(syllabification)理论认为音节是底层音段生成的,这等于不承认基本音节或底层音节形式。但是,母语者能够很容易判定基本音节的对立和音节的数量,判定音段、音位的对立和数量要困难得多,根据这一理由,我们认为音节是比音段更基本的概念,因此提取基本音节也比提取基本音段更重要。

就延展原则本身来说,其作用主要是断定音串(音段序列)中音节的数量和音核所在的位置,并不能断定音节边界,其他原则也不能。在没有核心音节以前,断定音节边界确实是一个复杂问题。但在提取语素并进一步提取核心音节的前提下,核心音节的边界基本上得到确认,核心音节在语流中可以有读音变化,会产生边界移动,但可以通过规则推导出来:

work+ing→wor-king
 sleep+ing→slee-ping
 take+ing→ta-king
 fall+ing→fa-ling
 feed+ing→fee-ding

以上可以概括出,在词中:

$$C_nVC + VC \rightarrow C_nV-CVC$$

从以上实例分析可以看出。如果不先提取语素音形,根据词和短语来提取音节,音节结构是五花八门的。汉语中可以找到这样的语素组合:

面包	唱啊
灵敏	行啊
蒙面人	等啊

从这些语素组合形式中可以提取出带 am、im、em 韵的音节,也可以提取 ŋa 这样的音节。但是这些都不是汉语的核心音节,这样的韵也不是汉语的核心韵。这些音节和韵都可以通过规则推导出来。

从这里也可以看出,从语素音形提取出来的语音单位,不一定总是协和的格局。有时候加上非核心音形,反而成了协和格局。比如,在不考虑非核心音节的情况,鼻音的组配方式是:

na	ma
an	aŋ

考虑非核心音节,组配方式成了:

na	ma	ŋa
an	am	aŋ

再从鼻音韵尾的配合看,m 在基本音系中不出现在韵尾位置上,如果加上 am、im、em,就可以出现在韵尾位置上,和其他鼻音韵

尾一致：

an	in	en
aŋ	iŋ	eŋ
am	im	em

可见,从语素音形提取的单位所组成的格局,通常是协和的,但这种趋势并不是绝对的。从语素音形提取语音单位的根本目的,是要获得不可由规则推导的核心语音单位。 ə 和鼻化的 ə̃ 也能够把 $[\text{ciər}^{55}]$ (心儿)和 $[\text{ciə̃r}^{55}]$ (星儿)区别开,但我们不把这种对立看成本音的对立,因为它们是在语素组合中发生的。

6.12 对立矩阵与音位的提取程序

在结构语言学中,音位被定义成最小的、有区别意义的语音单位。这里所说的区别意义包括了区别词的意义、区别话语的意义。前面我们已经讨论过,用词或词以上的组合不能有效提取核心语音单位,因此结构语言学的这个定义存在一定的问题。有了语素音形的前提,我们可以把音位定义成最短的、有区别语素音形作用的单位。到目前为止,音系分析是否需要音位这一层单位仍然是有争议的。这个争议本质上可以归结为是否承认音段对立的价值。Chomsky(1957)很明确的批评过对立的理论。

不过在语言调查中,有音位的概念是十分方便的。而音位负担的计算、音系协和等也要用到音位。我们有更充分的理由证明音位的存在价值(前面已有部分论述)。其中最重要的理由就在于音段记录的宽严两种方式都不可能绕开对立。比较下面实例:

铺 p'u^{55}	toŋ^{55} 东、冬
坡 p'o^{55}	tuŋ^{55} 东、冬

从严式记音看,“东、冬”中的元音可以有 o、u 不同的读音,并不区别语素音形,“铺、坡”中的元音读成 o 和 u 则区别语素音形,因此,在汉语中如果把 u 和 o 作为一个音位来处理,就不能说明在什么条件下读 u,什么条件下读 o,用语音规则来说,就不确定在什么条件下改写成 u,在什么条件改写成 o。当然,语素音形的对立不完全是音段对立造成的,比如“思”和“师”的对立。语素音形的对立不一定要归结为某个音段的对立,只要坚持严式记音,就能够把语素音形区别开。问题在于,“冬、东、懂、动、侗、洞、冻”等一系列字,会有不同的记音,把这些字的元音记录成两个音对母语者来说是没有价值的。更一般地说,任何一个严式记音,都可进一步记录成更严格的不同的音。这时候必须依赖是否对立来限制不必要的区分。另一方面,采用宽式记音,可以通过规则来预测音值,比如普通话 a 的三个读音。e 的几个读音也可以通过规则来预测。这样元音可以减少很多。问题在于,元音减少到多少就会影响规则? 比如上面的 u 和 o 显然是不能处理成一个音位的,或者说不能记录成相同的音段的。因此,音段的对立是音系中的一个必要概念。

下面的分析是在假定承认音位的情况下,通过音位提取程序来讨论音位的相对性和绝对性问题。在分析普通话音位系统时,人们似乎没有迫切感觉到需要一个音位提取程序,这是因为有充分的时间反复更正核对音位和音子的关系。在汉语方言的田野调查中,尤其是民族语言的田野调查中,需要在一定的时间内获得音位,如果没有充分的时间更正和核对材料,常常会出现问题。我们在田野调查时,对很多已经发表的材料和调查报告进行了核对,发现了不少问题。因此,一个有效的音位提取程序在田野工作中是必要的、迫切的。另一方面,如果我们展开音位提取程序的讨论,可以进一步认识到音位对立的本质。下面我们来讨论这样的程序。

6.12.1 对立研究的必要性

对立是音系研究的核心问题。这一思想是 Bloomfield (1933, 8.1) 强烈的主张, 他认为抛开意义是否相同来研究语音, 或者说纯语音研究, 是不重要的。我们同意 Bloomfield 的这种观点。不过 Bloomfield 以及很多学者只讨论音子的对立, 即音位的提取, 从我们前面的讨论看, 这种对立描写是不充分的。

从转换生成语法开始, 音位对立的思想开始受到批评。Chomsky (1957, 9.2.2, 9.2.3) 认为, 前人把语义区别作为话语在音位上对立的充分必要条件。Chomsky 用下面的证据批评了这种观点:

1. ration 的两种读法。语义相同, 语音不同。
2. bank 的两种意义。语义不同, 语音相同。

因此 Chomsky 不主张用区别意义来确定音位。这和 Harris (1951) 似乎是一致的。Chomsky 和 Halle 后来甚至主张抛弃音位 (Chomsky 和 Halle, 1968)。根本原因都在于怎样看待对立。

当 Chomsky 用 ration 的两种读法来证明区分音位可以不根据意义时, 我们也可以这样来回答问题: [æ] 和 [ei] 的对立在其他地方是存在的, 如 hate, hat。汉语的“亚”也是这种情况, 尽管 51 和 214 在“亚”这个字上不区别意义, 但 51 和 214 在其他语境下是区别意义的。因此可以这样来陈述问题: 如果两个音至少有一处对立, 本族人就能把两者区分成不同的音, 即使在不对立的情况下, 也能意识到两者读音不同。因此 Chomsky (1957) 的理由不能证明音位不重要。其实 Chomsky (1957) 关于 ration 的两种发音, 涉及本音和变音问题, ration 的两种读音实际上是在语素音形组合的过程中形成的, 是一种规则音变。

当然,语言中确实有很多音,发音差别是很大的,比如汉语普通话的 m 和 ŋ,并不对立。不过对立的思想并不要求不对立的音一定相似,只要求对立的音一定是不同的音。

对立是有层面的,可以有语音特征的对立、有音段的对立、语素音形的对立,汉语还可以有声母对立、韵母对立,等等。低层面的对立并不能充分解释高层面的对立,比如区别特征的对立系统并不能充分解释音段的对立系统。Halle(1959)关于俄语的分析让我们看到了语音特征在解释语音规则中的重要性,但不能证明音位不重要,也不是取消音位的理由。Halle(1959)分析了俄语中的下面现象:

$$[t] + [b] \rightarrow [d] + [b]$$

$$[p] + [g] \rightarrow [b] + [g]$$

认为这类现象可以用语音特征描写:

$$[\text{清}] + [\text{浊}] \rightarrow [\text{浊}] + [\text{浊}]$$

Halle 观察到音段可以还原到特征的现象。但正如我们前面分析的,不是所有的音段对立都可以还原成特征对立。当然,生成音系学家也可以说他们是用特征代替音位,并不是以对立的特征或者说区别特征代替音位。但事实上生成音系学的描写过程也不能避开对立问题,比如语素音形的对立。Halle (1959), Chomsky (1964, 1966b), Chomsky 和 Halle (1965), Postal (1962, 1968), Chomsky 和 Halle (1968) 等不承认音位层面 (phonemic level) 的存在,因此也没有音位 (phoneme) 和音位分析 (phonemic analysis)。Chomsky 和 Halle (1968) 取消音位和音位变体,直接用建立在音段基础上的音系表达 (phonological representation) 或底层音系表达 (underly-

ing phonological representation)通过语音规则生成语音输出。问题就在于 Chomsky 和 Halle 的音系表达是怎么得到的? 更进一步说,音段是怎么得到的。他们认为是从心智上构建出来的(mentally constructed, Chomsky and Halle, 1968, p14), 并认为传统的字母拼写(conventional orthography)是最佳的词汇表达式。这在方法论上存在问题。实际上传统的词的字母拼写也是以音位为基础的一种归纳,不过是不自觉的,因此也不是最佳拼写方案。当语言学家对目前没有文字的语言制定音位文字时,首先需要归纳音位系统。从这种意义上说, Chomsky 和 Halle(1968)以音段字母拼写为音系表达的基础仍然是以音位为音系表达的基础。

Chomsky 和 Halle(1968)关于音段的概念是没有严格定义过的概念,实际上已经含有音位的很多特点。记音过程有宽有严,这是不能避开的问题,但具体操作中已经含有音位处理的思想。把音段分得过细是没有必要的,比如把汉语 a 的三个变体都算成不同的音段就没有必要,这三个变体既不对立,音值区别也不大。另一方面,我们不可能把音段处理得太粗,普通话元音不能少于 5 个(详后)。把音子归纳成音位,就是在处理这个根本问题。

生成音系学取消音位的理由并不充分。前面我们已经分析过这样的实例:

tiŋ⁵⁵ 钉

təŋ⁵⁵ 灯

没有任何一种音系理论会把这两个语素描写成一样的读音,包括生成音系学,原因就在于这是两个对立的语素音形。对立是不能把这两个语素描写成同一个语素音形的关键。如果承认这个前提,立刻就可以把这两个语素音形的对立还原成 /i/ 和的

/ə/两个音子的对立。而前面我们已经讨论过,这两个音子的对立在汉语中不能还原成语音特征的对立。于是我们必须承认“钉”和“灯”的对立从根本上说是/i/和/ə/两个音子对立的结果,承认音子的对立就是承认音位。当然,承认音子对立并不是说音子不对立音节就一定不对立。如果音段 A 和音段 A' 不对立,音段 B 和音段 B' 不对立,语素音形 AB 和 A'B' 却可以对立。更一般地说:如果音段 A 和音段 A' 对立,以 A 为成分的序列 X1 和以 A' 为成分的语音序列 X2 一定对立。如果音段 A 和音段 A' 不对立,以 A 为成分的序列 X1 和以 A' 为成分的语音序列 X2 不一定不对立。从这种意义上看,音子对立是音子组合对立的充分条件,但不是必要条件。认识到音子的对立对预测语音序列的对立是有意义的。

一般地说,对于下面的包含关系:

特征	音段	语素音形
a	f1(a)	F1(f1(a))
b	f2(b)	F2(f2(b))

如果 a 和 b 对立,f1(a)和 f2(b)也对立,F1(f1(a))和 F2(f2(b))也对立。但反向的对立蕴含关系不成立,即 F1(f1(a))和 F2(f2(b))对立并不意味着 f1(a)和 f2(b)也对立,f1(a)和 f2(b)也并不意味着 a 和 b 也对立。也就是说,语音特征的对立并不能充分解释音段的对立,音段的对立并不能充分解释语素音形的对立,其他小于语素音形的语音片断也不能充分解释语素音形的对立。只有语素音形全部列出来,有多少对立的形式才能最终确定。语素音形的对立是最基本的对立。当然小于语素音形的语音片断或特征的对立仍然是有价值的。音位对立把语素音形对立还原成较小单位的对立,特征对立把语素音形对立还原成更小单位的对立,尽管这种还原不是处处成功的,但由于音位、区别特征在很多情况下能解释语音组合规则,所以音位、区别特

征是必要的。

更为重要的是,田野调查中没有绝对的准确记音或同质记音,因为没有绝对的准确发音或同质发音。实际记音中,会记录下各种复杂的音。发音人存在弥散发音现象,调查者存在弥散感知现象,这时判定同与不同必须靠对立。比较汉语普通话:

tiŋ ⁵⁵ 钉	tuŋ ⁵⁵ 东
təŋ ⁵⁵ 灯	toŋ ⁵⁵ 东

前面我们说/i/和/ə/是两个不同的音段,是因为形成了对立。右边一系列的/u/和/o/却只是一个相同的音段,因为没有对立。但在傣语中,[tuŋ]和[toŋ]两个音必须分开。对立在听音和记音过程中起到了根本的作用。

普通话声调从发音和听音两个角度都存在明显的弥散现象。单独读一个上声字,下面都是可能的:

214、213、212、22、11、112、113

无论弥散域怎样变化,上声一般不和阴平、阳平、去声混淆,要保持对立。因此,发音和记音可以有偏差,但不能超越对立,这是音系的根本性质。音段对立反映了母语者关注哪些对立,不关注那些对立。在语素音形的区别中,应该体现这种对立,这是用音位描写语素音形的另一个意义。

后结构语言学的双向唯一性原则曾受到生成音系学家批评,这是生成音系学取消音位的另一条理由。问题是音位理论并不必须遵守这条原则。音位存在的根本理由在于是否承认音段对立。

元音和谐或许也是考虑抛弃音位的理由之一。观察下面材料:

Turkish	English
evim	my house
ulusum	my nation
kolum	my arm
gülüm	my rose

这是 Clark(2000, p399) 在讨论韵律音系学的时候用到的一个材料, 现在我们用这个材料来考虑音位问题。通过对比不难发现, 土耳其语都有一个后缀, 相当于英语的 my。孤立地看 m 前的元音是不确定的, 但可以通过和前一元音和谐的规则确定。现在的问题是, 能否用元音和谐来否认音位对立的价值? 其实土耳其语也存在音位对立, 元音和谐和音位对立并不矛盾, 这个过程实际和汉语的轻声音节声调不确定相同。轻声音节也是没有自己独立的声调, 其具体读音在语流中由前面音节的声调决定。作为独立的语素 house, 土耳其语仍然有自己的读音, 这里的变化是在语素组合中形成的。这和汉语儿化变韵也没有根本区别。比如土耳其语语素 kedi(猫)、ev(房子)、ulus(民族)、kol(手臂)等, 都有自己的语素音形。

以上分析说明, 对立是音系研究的重要标准, 由此形成的各层面的对立单位都是有价值的。

现在我们考虑两种观点的统一。当我们进行田野调查时, 总有一个记音宽严的问题, 但一般不会把对立的音记录成相同的音, 这是在使用对立原则。至于不对立的音, 最严格的记音可以让每两个语素音形都没有相同的音段, 最宽的记音可以让不对立的位置都是相同的音段, 前提是两个不同语素音形至少有一个位置的音段是不同的。比如 [an] 和 [aŋ] 中的两个元音, 如果我们都记录成一个音, 实际上就在不自觉地归并音位, 如果记录成不同的音, 实际上就认为两个音不相似。音位理论的本质无非是把对立的音段分开(对立原则), 不对立的音段或者分开(认为不相

似),或者不分开(认为相似)。从这种意义上说,没有音位观念和不使用音位理论的人,是一种自然音位观,使用音位理论的人,是一种自觉音位观。自觉音位观无非指出了相似和不相似的理由。

在语言学中,认识到语音单位不完全是物理的,而是心理的,是一大进步,通过对立来进一步澄清心理印象的模糊概念,使判定语音单位有了可操作性,是另一个进步。Courtenay. J. N. (库尔德纳,1845—1929)最早认识到语音的功能或作用,认识到语音的生理性质、物理性质并不一致,认为语音是心理上的印象。我们知道,在生理上和物理性质上有明确音值的 /i/ 和 /y/ 两个音,在汉语普通话中是两个不同的语音单位,在汉语云南话中是一个单位。这一事实证明心理印象说是有道理的。Saussure(1912)关于语言是形式而不是实质,也是支持心理印象说的。Sapir(1921)中关于语音格局的思想也支持印象说。Saussure 还更明确地说语音的听觉与器官的发动形象一样直接,语音印象是整个音位理论的自然基础(Saussure, 1912, p67)。Saussure 还进一步指出对立的重要性(1912, p60):

我们即使把产生每个音响印象所必须的一切发音器官的动作都解释清楚,也没有阐明语言的问题。语言是一种以这些音响形象在心理上的对立为基础的系统。

也就是说,用对立来判定语音的异同的思想已经由 Saussure 首先提出。布拉格学派(1931)在《音位学术语标准化方案》中也提出了音位对立的观念。Bloomfield(1933)把通过对立提取音位的思想完全具体化了,提出了一套提取音位的操作手续。

对立是一种功能的观点,遭到过一些重要人物的反对。Chomsky 和 Halle(1968)抛弃了音位,本质上是抛弃音段对立的观念,同时引进了语音的深层形式。Chomsky 和 Halle 认为深层形式本质上是抽象形式,是为了描写复杂的表层形式而设立的。

抽象形式的思想 Hjelmslev, l. (叶尔姆斯列夫)也提到过,即音位是为了描写方便的抽象形式。现在的问题是,这种抽象性有没有一个认同的基础? Chomsky 和 Halle 没有讨论这个问题, Hjelmslev 也没有讨论过这个问题。我们认为这个认同基础就是对立,因为两个对立的音子不可能属于同一个深层形式或抽象形式,两个对立的语素音形也不可能属于同一个抽象语素形式,因此,深层形式或抽象形式的同一性判定不可能绕过对立的思维。

不仅抽象形式的理论不能绕开对立, D. Jones 主张具体物理属性的语音理论也不可能越过对立。鉴于心理印象或语音印象不容易把握, D. Jones 曾经从物理的角度把一个音位理解成特定语言中发音相关的一个集合,这个集合中的成员是相似的,分布是互补的。这实际上就是互补相似原则。Jones 试图从物理学的可观察角度来界定音位。问题是,这个集合中的成员如何确定,即相似到什么程度算是这个集合中的成员。实际上,两个对立的音段绝不可能理解成同一音位中的集合成员。也就是说,一个集合中的各个成员必须是不对立的,才能有资格理解成一个音位。两个音从物理的角度看无论有多相似,只要有对立,就不能作为这个集合的成员。因此,建立在音子集合观念上的音位也根本不可能超越对立。

6.12.2 元音音位的提取

赵元任《音位标音法的多能性》(1934)深入讨论了音位的相对性。赵元任提到的音位分析的相对性主要有两层含义。一是把两个互补分布的音素归纳成一个音位还是两个音位是相对的,二是线性方向的动态音切分成两个还是一个是相对的。下面我们在选定一种音段切分的基础上,重点讨论音素对立和互补的问题。赵元任所说的音素,下面我们用音子来代替,主要是考虑到音素和音标对应,是绝对的,而音子是特定语言的最短音段(也可以直接称为音段),已经有了调查者记音宽严的影响。

从赵元任的分析中容易看出,在把音子归纳成音位的过程中,有对立和互补相似两个根本原则,音位的相对性主要是因为不同的调查者对互补相似原则中相似程度的不同理解带来的,由此出现了归纳音位的不同方案。元音音位少的只有 3 个元音 Hartman(1944,p118),多的可以达到几十个不等。好像音位的归纳是相当随意的。不解决这个问题,尽管我们前面为音位作辩护,音位的地位仍然要受到质疑。

我们现在的核心问题是,音位有没有绝对性或确定性? 如果没有,我们认为会构成取消音位的最强硬的理由。

为了讨论方便,我们先从音子说起,并把问题限制在普通话元音的分析中。音子是一个语言中直接用国际音标记录下来的最短音段,不严格地说也可直接称为音段,比如汉语“良”中的 l、i、a、ŋ 都是音子。由于汉语元音音子都分布在韵母中,为了下面的分析更清楚,我们先列出韵母的严式记音。

[普通话韵母严式记音表]

	i[i],[ɿ]	u[u]	ü[y]
a[A]	ia[iA]	ua[uA]	
o[o]		uo	
e[ɤ]	ie[iɛ]		üe[yɛ]
ai[ai]		uai[uai]	
ei[ei]		uei[uei]	
ao[au]	iao[iau]		
ou[əu]	iou[iəu]		
an[an]	ian[iɛn]	uan[uən]	üan[yæn]
en[ən]	in[in]	uen[uən]	ün[yn]
ang[aŋ]	iang[iaŋ]	uang[uaŋ]	
eng[əŋ]	ing[iŋ]	ueng[uəŋ]	
ong[uŋ]	iong[iuŋ]		
er[ər]			

有人把[ər]看成是一个音段,理由是儿化的动作从发元音时就开始了,如果从这种角度出发,鼻音韵尾韵也是从发元音时就开始了有鼻音动作了。按照这种方式分析普通话音位会有另一种结果。在下面的分析中,我们暂时按照上面韵母表中一个音标一个音段的方式来处理问题,不影响讨论的实质。

前面提到过,两个音子在某一语音条件下不对立,在另外的语音条件下却可能对立,这意味着确定两个音子是否对立必须考虑全部可能出现的语音条件,只要在一处语音条件下是对立的,两个音子就是对立的。两个对立的音子绝不能够归并成相同的音位。另一方面,两个永远不对立的音子,应该归并到哪些音位中,也要考察全部音子来确定相似情况。因此,前面讨论的对立和互补仅仅是最基本的原则,这种原则要全面贯彻下去,需要对语素音形进行全部系统的分析。因此,音位分析的一个首要工作是提取音子对立信息,这包括两方面的内容:

1. 哪些音子和哪些音子对立;
2. 有多少对立项。

为了使对立项分析的过程更简单有效,下面给出一个我们在田野调查中经常使用的音子对立矩阵。音子对立矩阵用来系统表达音子对立项。以汉语普通话元音音子为例,下面普通话17元音音子对立矩阵的第1列和第1行是全部元音音子的有序排列,第1列的每个音子和第1行的每个音子都有一个相交的格,可以称为关系格,凡是对立的两个音子,都可以找出一个对立实例放在关系格中,凡是不对立的音子,在关系格中都留下空格,形成互补关系(为了节省空间,矩阵中对立实例省去了声调,事实上给出的实例都要求声调相同):

[汉语普通话 17 元音音子对立矩阵]

	ɿ	ʅ	i	ɨ	y	u	ʊ	e	ɛ	ə	o	ɐ	æ	a	ʌ	ɑ
ɿ	0					sɿ/su									sɿ/sʌ	
ʅ		0				ʂɿ/ʂu									ʂɿ/ʂʌ	
i			0		ly/li	lu/li				ən/in	po/pi			an/in	lʌ/li	ɑŋ/iŋ
ɨ				0						əŋ/iŋ						
y					0	ly/lu										
u	sɿ/su	ʂɿ/ʂu	lu/li		ly/lu	0				əŋ/uŋ	po/pu				lʌ/lu	ɑŋ/uŋ
ʊ							0									
e								0						ɛɪ/aɪ		
ɛ									0					ɪɛ/ɪa		
ə	sɿ/sʌ	ʂɿ/ʂʌ	lʌ/li		lʌ/ly	lʌ/lu				0				ən/an	lʌ/lʌ	
o			ən/in	əŋ/iŋ	ən/yn	əŋ/uŋ					0					əŋ/ɑŋ
ɐ			po/pi			po/pu									pʌ/po	
æ																
a			an/in		an/yn			ɛɪ/aɪ				ən/an		0		
ʌ	sɿ/sʌ	ʂɿ/ʂʌ	lʌ/li		lʌ/ly	lʌ/lu					pʌ/po				0	
ɑ			ɑŋ/iŋ			ɑŋ/uŋ				əŋ/ɑŋ						0

以上两个对立的音子可能在很多环境下对立,也可能只在一处对立,只要在一处对立,就是对立的音子,所以只列出了一处对立。

根据表头的音子计算,这个对立矩阵共有 17 个元音音子。如果元音音子 A 和元音音子 B 有对立,则 A 所在的行和 B 所在的列的相交格有对立的实例。计算下来共有 62 个对立实例。这个表是一个对称的表,以 0 构成的连线为对称中线,所以实际对立是 31 个。

把对称连线都对立的实例都列出来,是为了更好从行或列中看出一个元音的对立信息,因为无论是从横行的角度还是从纵列的角度,都容易看出一个音子可以和其他哪些音子构成对立。比如,在 17 个元音音子中,u 和其他元音音子产生的对立最多,共有 8 个:

	ɪ	ʊ	i	ɨ	y	u	ʊ	e
u	sɪ/su	ʂʊ/ʂu	lu/li		ly/lu	0		

	ɛ	ʏ	ə	o	ɛ	æ	a	ʌ	ɑ
u			əŋ/uŋ	po/pu				lʌ/lu	ɑŋ/uŋ

与此形成另一个极端,以上 17 个元音音子中,[u、ɛ、æ]所在的行和列全都没有对立项,全是空的:

	ɪ	ʊ	i	ɨ	y	u	ʊ	e	ɛ	ʏ	ə	o	ɛ	æ	a	ʌ	ɑ
u							0										
ɛ													0				
æ														0			

其他音子所在的行和列至少存在一项对立。

在对立矩阵中,凡是有空格的地方,空格所在的行和列的两个元音,都是互补的,存在着归并的可能。现在来考虑归并过程中的细节。如果把[u、ɛ、æ]所在的行和列都删除,17 个元音音子的对立矩阵就缩减为 14 个元音音子的对立矩阵:

[普通话 14 元音音位对立矩阵]

	1	ɿ	i	ɪ	y	u	e	ɛ	ɤ	ə	o	a	ʌ	ɑ
ɿ	0					sɿ/su			sɿ/sɤ				sɿ/sʌ	
ɿ		0				sɿ/ɤu			sɿ/ɤɤ				sɿ/ɤʌ	
i			0		ly/li	lu/li			lɤ/li	ən/in	po/pi	an/in	lʌ/li	ɑŋ/iŋ
ɪ				0						əŋ/iŋ				
y			ly/li		0	ly/lu			lɤ/ly	ən/yn		an/yn	lʌ/ly	
u[u]	sɿ/su	sɿ/ɤu	lu/li		ly/lu	0			lɤ/lu	əŋ/uŋ	po/pu		lʌ/lu	ɑŋ/uŋ
e						0						ɛɪ/aɪ		
ɛ								0				ɪɛ/ia		
ɤ	sɿ/sɤ	sɿ/ɤɤ	lɤ/li		lɤ/ly	lɤ/lu			0				lʌ/lɤ	
ə			ən/in	əŋ/iŋ	ən/yn	əŋ/uŋ				0		ən/an		əŋ/ɑŋ
o			po/pi			po/pu					0		pʌ/po	
a[æ,ɛ]			an/in		an/yn		ɛɪ/aɪ	ɪɛ/ia		ən/an		0		
ʌ	sɿ/sʌ	sɿ/ɤʌ	lʌ/li		lʌ/ly	lʌ/lu			lʌ/lɤ		pʌ/po		0	
ɑ			ɑŋ/iŋ			ɑŋ/uŋ				əŋ/ɑŋ				0

根据相似原则,删除的音子都归并在相似的其他音子中,相似的依据暂时根据普通话拼音方案。·由于进行了归并,说明表中的音段已经不只是音子,而是音位。这时的对立矩阵表不再称音子对立矩阵,而称音位对立矩阵。在 14 个元音音位对立矩阵中,每一个元音音位都可以和其他某个元音音位形成对立。不过这时的归并并没有影响到对立项的改变,缩减后的矩阵中,对立项仍然是 31 个,这是因为删除的音子是处处不对立的音子。我们把删除不对立音子后得到的音位对立矩阵称为元音音位初始对立矩阵。 $[u, \epsilon, \text{æ}]$ 的存在与否并没有影响到元音初始对立音位数目和初始对立项数目。这说明在记音过程中,有的音子的记音宽严并不影响元音音子对立项。能够体现对立的音子是记音过程中不能漏掉的第一个重要信息,初始对立矩阵就是反映这一信息的。

一般地说,一旦记音材料确定下来,所得到的元音音位初始对立矩阵是唯一的。从前面给出的普通话记音出发,普通话元音初始对立音位是 14,初始对立项是 31。

这种“唯一性”是建立在“记音材料确定下来”这一前提下的,不是建立在“记音准确”这一前提下的,因为从更严格的意义上说,记音的准确程度是相对的,如果我们把 uan 中的 a 记成 $[e]$,而 an 中的 a 仍然保留 $[a]$ 的记音,就会形成新的对立项 $[e]$ 和 $[a]$, $[uen/uən]$ 就是形成这种新的对立的语境,这时普通话元音初始对立矩阵的初始对立项会由 31 增加到 32,初始对立音位也会增加到 15。另一方面,如果不把 $[iŋ]$ 记录成 $[iŋ]$,而是直接就记录成 $[iŋ]$,就不会有 $[i/ɔ]$ 对立项,普通话初始对立矩阵的初始对立项就会由 31 减少到 30,不过初始对立音位仍然是 14。可见,有的音子记音的宽严会影响到初始对立音位和初始对立项,有的只影响到对立项的数目,不影响对立音位的数目。

上面矩阵中把“东”的韵母记录成 $[uŋ]$,则 $[ə]$ 和 $[o]$ 没有对立。

如果把“东”的韵母理解成[$o\eta$]而不是[$u\eta$],则[ə]和[u]在- η 前面不对立,[ə]和[o]在- η 前面对立,上面元音矩阵中[$\text{ə}\eta/u\eta$]的对立变成了[$\text{ə}\eta/o\eta$]的对立,这两种不同的记音都不会影响到初始对立音位和初始对立项,但改变了对立方式。这里仍然暂时把“东”的韵母记成 $u\eta$,后面的分析也是建立在这个基础上。

从总的趋势看,记音的宽严和对立音子以及音子对立项的多少是相关的,记音越严,音子对立项越多,对立音子也越多,记音越宽,音子对立项也越少,对立音子也越少。

初始对立矩阵中有很多空格,归并音位的实质就是合并对立矩阵中的音子,减少空格。现在我们考虑怎样在初始对立矩阵的基础上消除空格归并音位。为了讨论方便,我们尽量照顾普通话音位归并的传统,以便看出归并的本质。首先看[$1, \text{ɿ}, i, i$]几个元音,尽管这几个元音可以和其他元音对立,但相互之间不对立。由于它们相似,可以归并到[i]上来,[$1, \text{ɿ}, i$]和其他元音的对立也就需要[i]来体现。从上面的 14 元音对立矩阵中容易看出,凡是和[$1, \text{ɿ}, i$]对立的元音,正好和[i]也对立,这样就可以直接删除[$1, \text{ɿ}, i$]所在的行和列,在[i]后面加上变体。结果形成 11 元音音位对立矩阵(见下表)。

在 11 元音音位矩阵中,还存在空格,可以继续归并。低元音[a, A, α]相似并且不对立,可以把[A, α]归并到[a]中。与[$1, \text{ɿ}, i, i$]相互间的关系不同,从上面 11 元音矩阵中容易看出,和[A, α]对立的元音,[a]不一定与之对立,如[$l\text{A}/lu, l\text{A}/l\text{ʌ}, p\text{A}/po, \alpha\eta/u\eta$],删除[A, α]所在的行和列后,这些对立都需要体现出来。为了显示对立的具体信息,归并后保留下来的对立项仍然用音子表示,比如[a]和[o]的对立仍然用对立项[$p\text{A}/po$]。[a]和[u]的对立项可以是[$l\text{A}/lu$],也可以是[$\alpha\eta/u\eta$]。我们选择了[$l\text{A}/lu$]。于是可以得到普通话 9 元音音位对立矩阵。

[普通话 11 元音音位 24 对立项矩阵]

	i	y	u	e	ɛ	ɤ	ə	o	a	ʌ	ɑ
i[ɿ, ʅ, ɿ]	0	ly/li	lu/li			lx/li	ən/in	po/pi	an/in	lʌ/li	ɑŋ/iŋ
y	ly/li	0	ly/lu			lx/ly	ən/yn		an/yn	lʌ/ly	
u[u]	lu/li	ly/lu	0			lx/lu	əŋ/uŋ	po/pu		lʌ/lu	ɑŋ/uŋ
e				0					ɛl/aɪ		
ɛ					0				iɛ/ia		
ɤ	lx/li	lx/ly	lx/lu			0				lʌ/lɤ	
ə	ən/in	ən/yn	əŋ/uŋ				0		ən/an		əŋ/ɛə
o	po/pi		po/pu					0		pʌ/po	
a[æ, ɛ]	an/in	an/yn		ɛl/aɪ	iɛ/ia		ən/an		0		
ʌ	lʌ/li	lʌ/ly	lʌ/lu			lʌ/lɤ		pʌ/po		0	
ɑ	ɑŋ/iŋ		ɑŋ/uŋ				əŋ/ɛə				0

[普通话9元音音位19对立项矩阵]

	i	y	u	e	E	ɤ	ə	o	a
i[ɿ, ʅ, i]	0	ly/li	lu/li			lx/li	ən/in	po/pi	an/in
y	ly/li	0	ly/lu			lx/ly	ən/yn		an/yn
u[u]	lu/li	ly/lu	0			lx/lu	əŋ/uŋ	po/pu	lʌ/lu
e				0					ɛl/al
E					0				iE/ia
ɤ	lx/li	lx/ly	lx/lu			0			lʌ/lɤ
ə	ən/in	ən/yn	əŋ/uŋ				0		ən/an
o	po/pi		po/pu					0	pʌ/po
a[æ, ɛ, ʌ, ɑ]	an/in	an/yn	lʌ/lu	ɛl/al	iE/ia	lʌ/lɤ	ən/an	pʌ/po	0

归并到此,空格仍然存在。[e、E、ɤ、ə]相似不对立。按照汉语拼音归纳音位的方式,我们还可以把[E、ɤ、ə]继续归并成e,为了体现[E、ɤ、ə]和其他音段的对立,处理方式和[A、a]归并到[a]相同,于是有:

[普通话 6 元音音位 13 对立项矩阵]

	i	y	u	e	o	a
i[ɿ、ʅ、ɪ]	0	ly/li	lu/li	lɤ/li	po/pi	an/in
y	ly/li	0	ly/lu	lɤ/ly		an/yn
u[ʊ]	lu/li	ly/lu	0	lɤ/lu	po/pu	lA/lu
e[E、ɤ、ə]	lɤ/li	lɤ/ly	lɤ/lu	0		ei/ai
o	po/pi		po/pu		0	pA/po
a[æ、ɛ、A、ɑ]	an/in	an/yn	lA/lu	ei/ai	pA/po	0

这就是普通话 6 元音音位的情况。可以看出,6 元音音位矩阵中形成 13 对立项,中间还留有空格,还可以进行归并。如果认为 o 和 e 相似,可以继续把 o 归并到 e 中,于是有:

[5 元音音位 10 对立项矩阵(o 归并到 e)]

	i	y	u	e	a
i[ɿ、ʅ、ɪ]	0	ly/li	lu/li	lɤ/li	an/in
y	ly/li	0	ly/lu	lɤ/ly	an/yn
u[ʊ]	lu/li	ly/lu	0	lɤ/lu	lA/lu
e[E、ɤ、ə、o]	lɤ/li	lɤ/ly	lɤ/lu	0	ei/ai
a[æ、ɛ、A、ɑ]	an/in	an/yn	lA/lu	ei/ai	0

这里已经不再有空格,每一个音位都和另一个音位对立,不能再归并,因为这时任何一种归并都可能破坏对立原则,使对立的元音被归并到一起。可以说,普通话元音音位不能低于 5 个,或者说普通话最少需要 5 个元音音位。我们把普通话 5 个元音音位系统称为普通话元音必要音位系统。普通话必要音位系统的对立数是 5。

普通话 5 元音的必要音位系统和归并方式没有关系。比如,在 6 元音 13 对立矩阵中,o 不仅和 e[ɛ、ɤ、ə]不对立,和 y 也是不对立的。如果把 o 归并到 y 中,就会出现下面的 5 元音 10 对立项矩阵:

[5 元音音位 10 对立项矩阵(o 归并到 y)]

	i	y	u	e	a
i[ɪ、ɪ̯、i̥]	0	ly/li	lu/li	lɤ/li	an/in
y[o]	ly/li	0	ly/lu	lɤ/ly	an/yn
u[u]	lu/li	ly/lu	0	lɤ/lu	la/lu
e[ɛ、ɤ、ə]	lɤ/li	lɤ/ly	lɤ/lu	0	ei/ai
a[æ、ɛ̃、A、ɑ]	an/in	an/yn	la/lu	ei/ai	0

这时候也不再有空格,不能再归并。不过相比之下,o 和 e[ɛ、ɤ、ə]在一般人心目中更相似,而不是和 y 更相似,所以普通话终极音位系统选择 o 和 e[ɛ、ɤ、ə]归并更容易被人接受。

其实从一开始,我们也可以不按照普通话音位的归并方式,而采取其他可能的方式归并,到最后的結果是,普通话元音的必要音位系统都不会低于 5 个元音音位。限于篇幅,这里没有展开。

前面说过,记音的宽严会影响到初始对立矩阵,但记音的宽

严并不会影响到必要对立矩阵。比如,无论把把“英”的韵母记录成[iŋ]还是成[iŋ],如果要获得元音音位必要对立矩阵,由于[i]和[i]是互补的,都会被归并。

前面还说过,对音值的不同理解会影响到初始对立矩阵,由于音位是从对立矩阵中归纳出来的,对音值的不同理解会进一步影响到音位。如果把“东”的韵母理解成[oŋ],则[ə]和[o]在[əŋ]和[oŋ]上对立,在归并过程中会有这样一个阶段:

[6 元音音位 14 对立项矩阵]

	i	y	u	e	o	a
i[ɪ、ɪ、ɪ]	0	ly/li	lu/li		po/pi	an/in
y	ly/li	0	ly/lu			an/yn
u[ʊ]	lu/li	ly/lu	0		po/pu	lʌ/lu
e[ɛ、ɤ、ə]	lɤ/li	lɤ/ly	lɤ/lu	0	əŋ/oŋ	eɪ/aɪ
o	po/pi		po/pu	əŋ/oŋ	0	pʌ/po
a[æ、ɛ、ʌ、ɑ]	an/in	an/yn	lʌ/lu	eɪ/aɪ	pʌ/po	0

这时是 6 元音 14 对立项而不是 13 对立项。最后留下的空格是[o]和[y]的关系格。这两个音子是不相似的,所以在[əŋ]和[oŋ]对立的前提下,应该有 6 个音位。如果为了消除空格把[o]和[y]归并,音位数目也不会低于 5。可见,无论把“东”的韵母理解成[oŋ]还是[uŋ],元音必要对立音位最少不会低于 5 个。最少音位数是一个语言的音系的实证属性,不受研究者的影响。

由于音子的互补分布,音位归纳有相对性。如果强调互补音子不相似的方面,归纳出的音位数量就比较多。如果强调互补音子相似的方面,归纳出来的音位数量就比较少。这种相对

性并不等于说音位归纳是完全任意的。以上分析证明,如果忠实地记录北京话的元音,北京话的元音音位数量不得少于 5 个,因为至少有 5 个元音两两之间是相互对立的。

必要对立或终极对立的意义在于,至少需要 5 个元音音位,才能通过语音规则把其他元音推导出来。

需要解释的是,以上必要音位数的获得是以忠实音子的音值及其对立为前提的。如果不坚持这一点,就会出现一些特殊的归并。Hartman (1944, p118) 把 [i] 处理成 /ji/, 把 [u] 处理成 /wi/, 把 [y] (Hartman 写成 [ɥ]) 处理成 /jwi/。这种办法把音子本身的对立处理成音子组合的对立。Hartman 的目的是要满足 3 个半元音音位:

j w r

然后再建立一个不能自成音节的高元音音位 /i/, 通过半元音音位和高元音 /i/ 音位的组合,构成下面的音节:

ji 姨 wi 屋 ri 日 jwi 鱼 rwi 褥

[ɪ] 和 [ɨ] 也不单独设立音位,而是通过声母和 i 的组合体现出来, /i/ 和 /ts、tsh、s/ 组合时是 [ɪ], 和 /tɕ、tɕh、ʃ/ 组合时是 [ɨ], 这样北京话就只有一个高元音音位。在这样的高元音系统中,不仅 [ɪ]、[ɨ]、[i] 归并到 /i/ 中,连 /i/、/y/、/u/ 的对立也没有了, /y/、/u/ 和 /i/ 的对立通过 /i/ 和不同介音的组合体现出来。薛凤生 (1986) 也基本接受了 Hartman 只设立一个高元音 i 的方法,并提出根据明清以来“十三辙”的押韵作为一个高元音音位的重要依据。

根据我们后面要讨论的最宽直接对立环境,我们可以得到 5 个对立的元音。

li⁵¹ 立 lu⁵¹ 路 ly⁵¹ 驴 lɿ⁵¹ 勒 la⁵¹ 蜡

Hartman(1944)是 3 个基本元音。下面我们做一个比较:

5 个基本元音	i	u	y	ɿ	a
Hartman 的基本元音		wi	jwi	e	a

关键的区别在于, Hartman 引进了[w]和[j]。这种办法仍然能够把 5 个不同元音区别开,其中的要点是把两个单元音改变成两个复合元音,但弱点是音子和实际发音有距离。相同的音子作韵肯定押韵,但 i 和 wi 是不押韵的,所以把[u]理解成[wi]不对。[y]理解成[jwi]勉强说得过去。Hartman 的方法无法显示北京话中实际存在的元音对立,也不能解释现代北京话“屋、日、姨”并不押韵的事实,他所设立的音位本质上是抽象音位。我们认为,音系分析要尽可能忠于音子,才能反映出语言演变的结构情况。我们提出一个一致原则:音子和实际发音要尽可能一致。当然,我们还可以追问什么叫一致。押韵是一个标准。我们希望以后的研究还可以找到其他标准,尤其是实验语音学的标准。

6.12.3 辅音音位的提取

根据以上方式,也可以确定普通话辅音音子的对立项,下面普通话辅音音子对立矩阵各个方格中的两个字都是对立的实例。这个矩阵也可以完全用语素音形表示,其中两个对立的辅音相交的方格中的韵母是对立出现的语境,空格的地方是没有对立的语境(如表所示)。

以上共 157(210-12-20-21)项对立。舌面音和舌尖、舌尖后、舌根三组的关系格都是空格,可以归入三组中的一组。如果归入舌根音,形成 270 页的 19 辅音对立矩阵表。

[普通话辅音子对立矩阵]

	p	p'	m	f	ts	ts'	s	t	t'	n	l	te	te'	ɕ	tɕ	tɕ'	ʂ	ʂ'	ʐ	k	k'	x	ŋ
p	0																						
p'	盼半	0																					
m	慢半	满盼	0																				
f	饭半	饭盼	饭慢	0																			
ts	暂半	暂盼	暂慢	暂饭	0																		
ts'	灿半	灿盼	灿慢	灿饭	灿暂	0																	
s	三班	三攀	伞满	三翻	伞攒	三掺	0																
t	淡半	淡盼	淡慢	淡饭	淡暂	淡灿	单三	0															
t'	探半	探盼	探满	探饭	探暂	探灿	贪三	探淡	0														
n	男盘	男盘	男蛮	男凡	男咱	男蚕	报伞	难蛋	男谈	0													
l	烂半	烂盼	烂慢	烂饭	烂暂	烂灿	懒伞	烂淡	蓝谈	蓝男	0												
te	见变	见骗	见面					见店	尖天	见念	见链	0											
te'	欠变	欠骗	欠面					欠店	千天	欠念	欠链	欠见	0										
ɕ	线变	线骗	线面					线店	先天	线念	线链	线见	线欠	0									
tɕ	站半	站盼	站慢	站饭	站暂	站灿	盏伞	站淡	站探	住怒	站烂			0									
tɕ'	颤半	颤盼	颤慢	颤饭	颤暂	颤惨	铲伞	颤淡	颤探	焯男	颤烂				颤站	0							
ʂ	善半	善盼	善慢	善饭	善暂	善惨	闪伞	善淡	善探	树怒	善烂				善站	善颤	0						
ʐ	染板	燃盘	染满	染反	染攒	染惨	染伞	染胆	燃谈	燃男	燃栏				染站	染铲	染闪	0					
k	肝班	肝攀	敢满	肝翻	敢攒	敢惨	敢伞	敢胆	肝贪	干难	敢懒				敢站	敢铲	敢闪	敢染	0				
k'	看半	看盼	砍满	看饭	看暂	砍惨	砍伞	砍胆	看探	看难	砍懒				砍站	砍铲	砍闪	砍染	砍敢	0			
x	汉半	汉盼	喊满	汉饭	汉暂	喊惨	喊伞	喊胆	汉探	含男	喊懒				喊站	喊铲	喊闪	喊染	喊敢	喊砍	0		

[普通话辅音子对立矩阵(韵母语境)]

[illegible]

[普通话 19 辅音对立矩阵]

	p	p'	m	f	ts	ts'	s	t	t'	n	l	ts	ts'	ʒ	k	k'	x	ŋ
p	0																	
p'	盼半	0																
m	慢半	满盼	0															
f	饭半	饭盼	饭慢	0														
ts	暂半	暂盼	暂慢	暂饭	0													
ts'	灿半	灿盼	灿慢	灿饭	灿暂	0												
s	三班	三攀	伞满	三翻	伞攒	三掺	0											
t	淡半	淡盼	淡慢	淡饭	淡暂	淡灿	单三	0										
t'	探半	探盼	探满	探饭	探暂	探暂	贪三	探淡	0									
n	男盘	男盘	男蛮	男凡	男咱	男蚕	报伞	难蛋	男谈	0								
l	烂半	烂盼	烂慢	烂饭	烂暂	烂灿	懒伞	烂淡	蓝谈	蓝男	0							
ts	站半	站盼	站慢	站饭	站暂	站灿	盏伞	站淡	站探	住怒	站烂	0						
ts'	颤半	颤盼	颤慢	颤饭	颤暂	颤惨	铲伞	颤淡	颤探	蝉男	颤烂	颤站	0					
s	善半	善盼	善慢	善饭	善暂	闪惨	闪伞	善淡	善探	树怒	善烂	善站	善颤	0				
ʒ	染板	燃盘	染满	染反	染攒	染掺	染伞	染胆	燃谈	燃男	燃烂	染盏	染铲	染闪	0			
k(tc)	肝班	肝攀	敢满	肝翻	敢攒	敢惨	敢伞	敢胆	肝贪	固怒	敢懒	敢盏	敢铲	敢闪	敢染	0		
k'(tc')	看半	看盼	砍满	看饭	看暂	砍惨	砍伞	砍胆	看探	库怒	砍懒	砍盏	砍铲	砍闪	砍染	砍敢	0	
x(c)	汉半	汉盼	喊满	汉饭	汉暂	喊惨	喊伞	喊胆	汉探	含男	喊懒	喊盏	喊铲	喊闪	喊染	喊敢	喊砍	0
ŋ										敢港								

[ŋ]除了和[n]在韵尾位置上对立,和其他辅音音子不对立,但也不相似,可以不再归并,于是普通话需要 19 个辅音音位。如果把 ŋ 归并到其他辅音中,比如归并到[m],这时矩阵中不再有空格,对立项目不能再减少,得到 18 个两两对立的辅音音位。可见普通话辅音音位不能少于 18。由于出现在音节开头的辅音也不能少于 18,因此普通话声母不能少于 18。

6.12.4 音子组合数与对立环境

前面说汉语普通话必要元音音位不应该少于 5,现在的问题是哪 5 个音子最有希望作代表元音。现在来考虑音子组合数和对立环境问题。

Hockett (1958, 2.1, p19)曾经用行(row)和列(column)的方式来考察音位对立,是一项系统的工作。但 Hockett 没有穷尽对立条件,因为没有首先提取语素音形,也没有从音子组合的角度考虑问题,使有些现象没有观察到。为了有效反映音子对立信息,我们先给出一个直接对立环境的概念。考虑下面的实例:

təŋ ⁵⁵	登
taŋ ⁵⁵	当

[t_ŋ⁵⁵]是把[ə]和[a]区别开的一个语音对立环境,这个对立环境可以进一步归约为-ŋ这一语音环境,这个语音环境不能再归约,这就是直接对立环境。严格地说,音子的直接对立环境和音子互为直接成分。下面把元音音子在直接对立环境中的对立排列出来:

续表

	i	u	y	a	A	ɑ	o	ɤ	ə	e	ɛ	ɪ	ɿ
s-		su			sA	.		sɤ				sɪ	
tʂ-		tʂu			tʂA			tʂɤ					tʂɿ
tʂ'-		tʂ'u			tʂ'A			tʂ'ɤ					tʂ'ɿ
ʂ-		ʂu			ʂA			ʂɤ					ʂɿ
ʐ-		ʐu						ʐɤ					ʐɿ
tɕ-	tɕi		tɕy										
tɕ'-	tɕ'i		tɕ'y										
ɕ-	ɕi		ɕy										
k-		ku			kA			kɤ					
k'-		k'u			k'A			k'ɤ					
x-		xu			xA			xɤ					
ʔ-	ʔi	ʔu	ʔy		ʔA			ʔɤ					

续表

	i	u	y	a	ʌ	ɑ	o	ɤ	ə	e	ɛ	ɿ	ʅ
-aŋ	iaŋ	uaŋ											
-ən		uən											
-əŋ		uəŋ											
-ɛ	ie		ye										
a-	ai												
ɑ-		au											
ə	əi	əu											
p-	pi	pu			pʌ		po						
p'-	p'i	p'u			p'ʌ		p'o						
m-	mi	mu			mʌ		mo						
f-		fu			fʌ		fo						
t-	ti	tu			tʌ					tɤ			
t'-	t'i	t'u			t'ʌ					t'ɤ			
n-	ni	nu	ny		nʌ					nɤ			5
l-	li	lu	ly		lʌ					lɤ			5
ts-		tsu			tʂʌ					tʂɤ		tsɿ	
ts'-		ts'u			ts'ʌ					ts'ɤ		ts'ɿ	

以上第 1 列是元音音子可能出现的直接对立环境。从第 2 列开始, 每一列可以反映出每一个音子能够出现在多少直接对立环境中, 我们把这个数量称为音子的组合数。矩阵中的第 2 行给出了每个音子的组合数(没有计算声调)。

以上矩阵的每一行都表示相同的直接对立环境中可以出现多少对立音子, 比如在 -n 这个直接对立环境中可以出现 [i、y、a、ə] 4 个对立音子。有少数行的最小对立环境不能保证有相同的声调, 比如 ts- 所在的行, 可以出现 4 个对立音子 [u、A、ɤ、ɿ], 但并不出现在相同的声调中:

阴平	阳平	上声	去声
tsu ⁵⁵ 租	tsA ⁵⁵ 扎		tsɿ ⁵⁵ 资
tsu ³⁵ 族	tsA ³⁵ 砸	tsɤ ³⁵ 泽	
tsu ²¹⁴ 组	tsA ²¹⁴ 咋(方言)		tsɿ ²¹⁴ 子
		tsɤ ⁵¹ 仄	tsɿ ⁵¹ 字

因此, 直接对立环境可以分成严式和宽式两种, 严式要求声调也相同, 宽式不要求声调相同。

?-和 l- 这两个环境中都可以出现 5 个对立音子, 并且都可以都限制在同一个声调:

i ⁵⁵ 一	u ⁵⁵ 乌	y ⁵⁵ 迂	ɤ ⁵⁵ 阿	A ⁵⁵ 阿
li ⁵¹ 立	lu ⁵¹ 路	ly ⁵¹ 绿	lɤ ⁵¹ 勒	lA ⁵¹ 蜡

其他对立环境都没有超过 5 个。我们说 ?-和 l- 是元音音子的两个最宽直接对立环境, 最宽直接对立环境也证明, 普通话元音音子至少有 5 个是两两对立的, 因此普通话元音音位至少不能低于 5 个。容易看出, 通过直接对立环境的排列也能够找出一种语言的最少音位数, 不过这种方法不能反映一个语言的音子到

底有多少对立对儿。要确定对立对儿仍然需要音子对立矩阵。

在音位归并中,音子组合数越高,越有理由充当音位的典型变体或底层音系表达式的音段,这样可以使底层表达式更有经验依据。在上面的元音音子最小对立环境矩阵中,有一个值得注意的现象:出现在最宽对立环境中的音子[i、y、u、A、ɤ],也正是音子组合数最高的前5个音子。这更有理由让我们把典型变体和组合指数高的音子联系起来。

如果只考虑元音音子在韵或韵母的中的组合情况,可以得到另一种组合指数,矩阵中的第3行列出了每个音子的在韵母范围内的组合指数(不包括单独作韵母)。

从以上对比可以看出, 1^{-51} 是完全对立条件,而且是严式的。在这种情况下,完全对立音子是5个。即使宽式对比,完全对立音子也不会多于5。

如果把对比放宽到非核心音节,比如考虑叹词“噢”,最宽对立环境?中出现的对立音子可以是6:

i⁵⁵ 一 u⁵⁵ 乌 i⁵⁵ 迂 ɤ⁵⁵ 阿 A⁵⁵ 阿 o⁵⁵ 噢

根据最宽对立环境,可以确定一个语言的最少音位数目。这个数据是本族人关于音位直觉的重要部分。一种语言最多可以有多少音位?由于音位提取有互补相似原则,音位最大数目可以不受对立的限制,最多可以多到和音子的数目相同。但最少音位数却是唯一的。

通过以上方法,可以找到声母的最宽对立条件(见下表)。以上辅音的最宽直接对立环境是18,这还是在不限声调的情况下的对立环境。如果限制声调,最宽直接对立环境数要少一些。比如上面对立环境数为18的,如果限制声调,都不足18:

[汉语辅音声母音子直接对立环境(组合指数)(不考虑声调)]

	p	p'	m	f	ts	ts'	s	t	t'	n	l	tɕ	tɕ'	ɕ	tʂ	tʂ'	ʂ	ʐ	k	k'	x
3					ɿ	ɿ	ɿ														
4															ɿ	ɿ	ɿ	ɿ			
10	i	i	i					i	i	i	i	i	i	i							
18	u	u	u	u	u	u	u	u	u	u	u				u	u	u	u	u	u	u
5										y	y	y	y	y							
17	ʌ	ʌ	ʌ	ʌ	ʌ	ʌ	ʌ	ʌ	ʌ	ʌ	ʌ				ʌ	ʌ	ʌ	ʌ	ʌ	ʌ	ʌ
4	o	o	o	o																	
14					uo	uo	uo	uo	uo	uo	uo				uo	uo	uo	uo	uo	uo	uo
14					ɤ	ɤ	ɤ	ɤ	ɤ	ɤ	ɤ				ɤ	ɤ	ɤ	ɤ	ɤ	ɤ	ɤ
16	ai	ai	ai		ai	ai	ai	ai	ai	ai	ai				ai	ai	ai	ai	ai	ai	ai
6															uai	uai	uai	uai	uai	uai	uai
14	ɐi	ɐi	ɐi	ɐi	ɐi	ɐi		ɐi	ɐi	ɐi	ɐi				ɐi	ɐi	ɐi	ɐi	ɐi	ɐi	ɐi
13					uei	uei	uei	uei	uei	uei	uei				uei	uei	uei	uei	uei	uei	uei
17	au	au	au		au	au	au	au	au	au	au				au	au	au	au	au	au	au
7	iau	iau	iau					iau	iau	iau	iau										
17		əu	əu	əu	əu	əu	əu	əu	əu	əu ^②	əu				əu	əu	əu	əu	əu	əu	əu
4			ieu									ieu	ieu	ieu							
18	an	an	an	an	an	an	an	an	an	an	an				an	an	an	an	an	an	an
10	ian	ian	ian					ian	ian	ian	ian	ian	ian	ian							

续表

	p	p'	m	f	ts	ts'	s	t	t'	n	l	tɕ	tɕ'	ɕ	tʂ	ʂ	ʐ	k	k'	x
14					uan	uan	uan	uan	uan	uan	uan				uan	uan	uan	uan	uan	uan
3												yan	yan	yan						
18	aŋ	aŋ	aŋ	aŋ	aŋ	aŋ	aŋ	aŋ	aŋ	aŋ	aŋ				aŋ	aŋ	aŋ	aŋ	aŋ	aŋ
5										iaŋ	iaŋ	iaŋ	iaŋ	iaŋ						
6																uaŋ	uaŋ	uaŋ	uaŋ	uaŋ
16	ən	ən	ən	ən	ən	ən	ən	ən	ən	ən					ən	ən	ən	ən	ən	ən
13					uən	uən	uən	uən	uən		uən				uən	uən	uən	uən	uən	uən
14					uŋ	uŋ	uŋ	uŋ	uŋ	uŋ	uŋ				uŋ	uŋ	uŋ	uŋ	uŋ	uŋ
3												iuŋ	iuŋ	iuŋ						
18	əŋ	əŋ	əŋ	əŋ	əŋ	əŋ	əŋ	əŋ	əŋ	əŋ	əŋ				əŋ	əŋ	əŋ	əŋ	əŋ	əŋ
10	iŋ	iŋ	iŋ					iŋ	iŋ	iŋ	iŋ	iŋ	iŋ	iŋ						
8	ie	ie	ie							ie	ie	ie	ie	ie						
3												ye	ye	ye						
8	in	in	in							in	in	in	in	in						
3												yn	yn	yn						

① 勃

② 癖

[汉语辅音音子组合指数(考虑声调)]

	p	p'	m	f	ts	ts'	s	t	t'	n	l	te	te'	e	ts	ts'	s	k	k'	x
-u ⁵⁵		铺		夫	租	粗	苏	都	涂						猪	出	书	孤	哭	呼
-u ³⁵	醅	仆	模	福	足		俗	毒	困	奴	炉				竹	锄	熟			湖
-u ²¹⁴	补	谱	母	辅	组			堵	土	努	鲁				主	楚	鼠	古	苦	虎
-u ⁵¹	布	堡	墓	赴			素	度	兔	怒	路				注	处	数	固	数	户
aŋ ⁵⁵	帮	兵		方	葬	仓	桑	当	汤	囊	啞				张	昌	伤	钢	康	夯
aŋ ³⁵		旁	忙	房		藏			唐	囊	狼					长			扛	行
aŋ ²¹⁴	绑	榜	蟒	仿	狙		噪	党	躺	囊	朗				长	厂	赏	港		
aŋ ⁵¹	棒	胖		放	葬		丧	荡	烫	囊	浪				丈	唱	尚	扛	炕	巷
an ⁵⁵ ①	班	潘		翻	簪	参	三	单	滩	团					詹	掺	山	甘	刊	罕
an ³⁵		盘	蛮	烦	咱	蚕			坛	难	兰					缠				含
an ²¹⁴	版		满	反	攒	惨	伞	胆	坦	赅	览				斩	产	闪	敢	砍	喊
an ⁵¹	半	判	慢	饭	赞	粲	散	蛋	探	难	滥				站	颤	扇	干	看	汗
aŋ ⁵⁵	崩	烹	懵	锋	增	噌	僧	灯							征	称	声	耕	坑	哼
aŋ ³⁵	甬	彭	盟	缝		层			藤	能	棱					成	绳			横
aŋ ²¹⁴	绷	捧	猛	讽				等			冷				整	惩	省	耿		
aŋ ⁵¹	泵	碰	梦	凤	憎	蹭		瞪			楞				邳	秤	圣	更		

① “慢慢”中的后一个音节是派生音节,没有计算。

以上最大对立环境是 17, 所以更严格地说, 最大对立环境只有 17。

对比还可以不断进行下去, 最宽直接对立条件是 [a、u] 等, 在这些条件下至少可以构成 18 个辅音对立。因此普通话辅音音位至少不少于 18。这项工作在田野调查中可以通过数据库很快完成。除去最宽直接对立环境中的辅音, 剩下的是舌面音和舌根鼻音。舌面音和舌根鼻音由于在 [i] 条件下和唇音、舌头音对立, 不能归并在这些音里, 因此只能归并在这些音以外的其他音里, 这时需要根据相似标准确定它们的归类。

思考问题

1. 音系分析程序的步骤是什么?
2. 核心语音单位能否独立于语音规则而获得?

7. 语义组合关系的线性分析

7.1 组合规则的语义解释

Chomsky 的 *Syntactic Structure* (1957)、*Three models for the description of language* (1956)、*The logical structure of linguistic theory* (1955) 等著作基本上奠定了形式语法分析的基础,或者说解决了形式组合的问题。数学中的形式文法理论也是在这种基础上展开的。从 1957 年到 1965 年,形式文法理论和转换理论的技术问题得到了详细的研究,除了乔姆斯基的工作外,J. A. Fodor、J. J. Katz、C. J. Fillmore、P. M. Postal, 都做了很多重要工作。

但在自然语言中,单位的组合不仅受形式规则的限制,而且受语义规则的限制。这是形式语法规则还不能解决的问题。从语义组合关系看,更需要的是—种限制,需要的是既符合语法(形式的限制)、又符合可接受性(语义规则限制)的句子。不考虑意义的形式规则生成能力很强,但生成能力强不见得是自然语言语法研究的目的。自然语言语法研究的目的是要既能生成所有可能的句子,又要防止生成不符合语法的句子。比如根据 Chomsky(1957)的规则,我们可以有:

$$S \rightarrow NP + VP$$

$$VP \rightarrow V + NP$$

可以得到这样一个抽象结构:

NP1 + V + NP2

并形成下面的句子:

汪锋吃了苹果

· 苹果吃了汪锋

从语义的角度看, NP1 和 NP2 尽管都是 NP, 但语义性质是不一样的。这就是说, NP1 必须是有动物特征的, 满足了这样的语义条件, NP1 + V + NP2 才是符合语法的。

在具体工作中, 语义限制条件的问题也日益突出。Chomsky 的一些学生率先展开了语义研究。比如最初 Katz 和 Fodor (1963, p170—210) 在语义研究中提出一个著名的公式:

语言描写 - 语法学 = 语义学

这个公式从某个角度看可以理解成语法研究不需要语义学的知识。如果我们把公式变换一下就有:

语言描写 = 语法学 + 语义学

即语言描写必须包括语义学。后来 Katz 和 Postal (1964) 又提出语法学中必须包括语义, 即语法也不能独立于语义学。正是在这一基础上, Chomsky (1965) 给出了句法、音系、语义三分的标准理论 (Standard Theory), 其中语义部分的主要观点就源于 Katz 和 Postal (1964) 的工作。Chomsky (1965) 标准理论的结构框架可以概括如下:

句法部分 (syntactic component)

基础部分(base component)

范畴部分(categorial component)

词库(lexicon)

(基础部分生成深层结构(deep structure))

转换部分(transformational component)

(转换部分在深层结构基础上映射(map)出表层结构(surface structure))

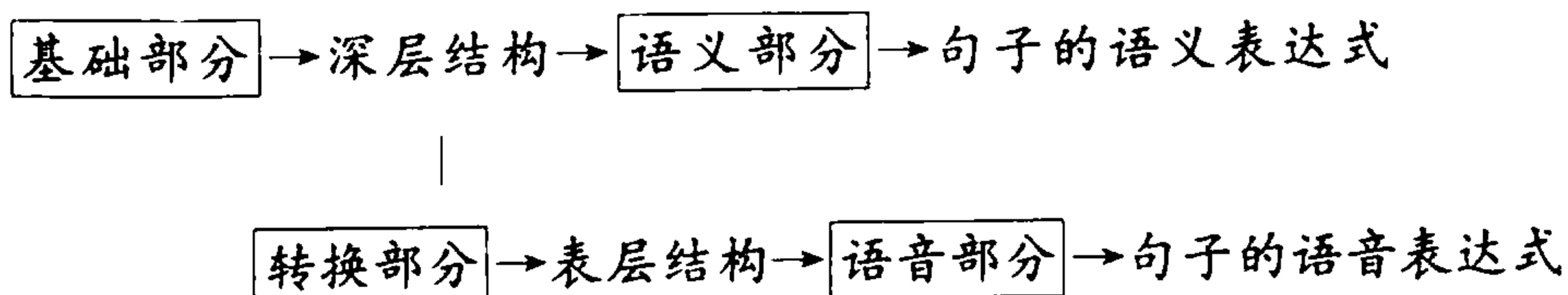
语义部分(semantic component)

(深层结构在这一部分得到语义解释,形成语义表达式)

音系部分(phonological component)

(表层结构在这一部分得到语音解释,形成语音表达式)

语法由句法部分、语义部分和语音部分组成。句法部分由一个基础部分和一个转换部分组成。基础部分包括范畴部分(短语结构规则)和词汇部分,这两个部分相互作用生成深层结构,深层结构在语义部分得到解释,形成语义表达式。转换部分则把深层结构转化为表层结构,表层结构经过音系规则得到语音表达式。各部分关系如下所示:



在深层结构解释语义,是标准理论的重要思想之一。

7.2 次范畴化、词库和平行周遍条件

从前面分析可以看出,语类的线性组合方式是组合规则的一个主要部分,但是仅仅通过 Harris 和 Chomsky 的语类线性组合规则来解释所有语符的线性组合的条件是不够的。所遇到的问题表现在两个方面。首先,正如前面已经提到的,符合语类规则的组合并不一定是可以接受的组合。以 VP 为例:

$VP \rightarrow V + NP$

杀鸡、杀鸭、杀牛

· 杀树、杀草、杀花

· 杀空气、杀水、杀纸、杀桌子

“杀树”一类组合不成立不是语言系统内部的组合规律问题,在有些语言中可以说。“杀空气”一类组合不成立是认知问题,所有的语言大概都不能这么说。当然这些不能说的组合也可以放在语义组合层面来解释,从语法层面上说并没有违反语类规则。不过语类规则遇到的另一个问题就不完全是语义问题或认知问题,跟早期语类规则在描写上的不充分性有关系。早期语类规则是二元的,而且过于简单化,即一个组合的语类是由两个直接成分的语类解释的,以 VP 为例,概括成改写规则就是:

$VP \rightarrow V + NP$

但是,二元语类推导模式并不充分。考虑下面模式:

送学生一本书($V + NP1 + NP2$)

教学生学英语($V + NP + VP$)

告诉学生美国历史很短(V+NP+S)

说出来一个错误句子(V+VP+NP)

还有更多的情况。用两个直接成分的语类来解释以上组合显然有困难。

Chomsky(1965)提出了进一步解决语类组合条件的基本方式。Chomsky(1965, 2.3)试图在语法中处理语义并将语义形式化。基本方法是次范畴化(subcategorization),以此来限制太强的生成能力,以生成合法的句子。下面我们来分析次范畴化,并尝试给出一个判定标准。

7.2.1 名词的次范畴化和动词的选择规则

Chomsky(1965)解决语义问题的办法是,先把名词分析成带有一束语义或句法特征的复合符号(complex symbol)。考虑下面两个实例:

sincerity may frighten the boy

* the boy may frighten the sincerity

为了描写上面句子的生成机制,他首先给出了名词的次范畴规则,比如对 sincerity 和 boy 的次范畴化过程是:

sincerity[+N, -可数的, +抽象的]

boy[+N, +可数的, +普通的, +动物的, +人类的]

这是将名词句法特征化,方括号中表示所含特征的具体情况,负号指没有这个句法特征,正号指有这个句法特征。sincerity 被分析成了带有一组句法特征(syntactic features)的复合符号,“-可数的”表示 sincerity 没有“可数”特征。同样的,boy 也被

分析成了带有一组句法特征 (syntactic features) 的复合符号。于是 sincerity may frighten the boy 这个句子中的 NP 如果用复合符号来表示就是:

[+N, -可数的, +抽象的] + M + V + the + [+N, +可数的, +普通的, +动物的, +人类的]

这里动词 V 前面的 M 指动词的情态、时态等变化部分。能够满足复合特征的名词进入句子后就是合法的句子, 否则就是不合法的句子:

sincerity may frighten the girl
sincerity may frighten the man
sympathy may frighten the boy
carefulness may frighten the boy

- * the boy may frighten the sincerity
- * the girl may frighten the sincerity
- * the boy may frighten the sympathy
- * the boy may frighten the carefulness

复合符号实际上把名词分成了很多小类, 因为只有满足复合符号条件的名词才能进入相应的句子, 这些名词自然就构成了一个小类。容易看出, 名词的次范畴化不是任意的引入句法特征, 而是为了让动词选择合适的名词而给名词引入句法特征, 所以名词的次范畴化必须参考动词选择名词的要求。于是 Chomsky 在名词次范畴化的前提下又讨论了动词的选择规则:

frighten: [+V, +__ NP(+人类的),]

名词的次范畴规则把名词分成更小的类,动词 frighten 在选择后面的名词时,只能选择有“人类”特征的 boy,不能选择没有“人类”特征的 sincerity。动词的 frighten 的这一特点也决定了名词次范畴化必须考虑“人类”这个句法特征。

名词的次范畴化和动词的选择条件是相互依赖的语义组合行为,因为一旦 frighten 前后 NP 的句法特征确定下来, frighten 位置上的动词小类也确定下来了:

sincerity may frighten the girl
sincerity may interest the girl
sincerity may horrify the girl
sincerity may worry the girl
sincerity may touch the girl
.....

和结构语言学的分类进行比较可以看出,次范畴化也是一种分类,但主要是从句法特征的角度而不是分布的角度给词分类。当然,boy 和 sincerity 在小类上的区别也可以从分布上得到解释。不过 Chomsky 不是从分布上来进行次范畴化或分小类,而是从语义特征“人类、抽象、可数、动物、普通”等方面展开次范畴化。这两种分类,即根据特征的分类和根据分布的分类,到底有没有区别?我们认为是有的。至少在分小类上,根据特征的分类有利于组合规则的描写。让我们先弄清楚分布分类的实质,如果我们从分布上进行罗列:

to frighten the boy
to frighten the girl
* to frighten the sinserity
* to frighten the book
.....

可以概括出两类 N:

Na: boy, girl, man, mother, father ...

Nb: sincerity, carefulness, sympathy...

其中 Na 类名词能出现在 frighten 的后面,以 Na 为核心的 NP 也能出现在 frighten 的后面:

to frighten the boy

to frighten the little boy

to frighten the little boy in the class

但是,对于那些不以 Na 为核心的 NP,是否可以进入 frighten 后面只能取决于是否具有“人类”的特征:

to frighten those that I learnt money from

to frighten those that I spoke to

* to frighten those that I bought yesterday

* to frighten those that I built before

再考虑汉语的“恐吓”一词,根据 N 是否出现在其后也可以把 N 概括成两个分布类:

Na: 孩子、妈妈、老师、学生、富人……

Nb: 桌子、房子、空气、水、火……

只有 Na 和以 Na 为核心的 NP 可以出现在“恐吓”后面:

恐吓孩子、恐吓妈妈、恐吓老师、恐吓学生、恐吓富人……

恐吓幼儿园的孩子、恐吓有钱的学生

* 恐吓桌子、* 恐吓房子、* 恐吓空气、* 恐吓水、* 恐吓火

但是,对那些不以 Na 为核心的 NP,是否可以进入“恐吓”后面,只能取决于是否具有“人类”的特征:

恐吓我不认识的、恐吓有钱的、恐吓胆小的

* 恐吓摩托罗拉制造的、* 恐吓含有三聚氰胺的

因此,根据句法特征或语义特征进行分类或确定组合条件,有时候是必要的。也就是说,有一些组合如果最终不给出一个特征来加以区别,就不能上升到规则,只能停留在记忆。可以说,Chomsky 的次范畴化,或早些时候一些语言学家涉及的次范畴方法,反映了人类语言能力的一个方面。

7.2.2 动词次范畴化

Harris(1946)的语类生成模型和 Chomsky(1957)的短语规则生成能力太强。名词的次范畴化规则是在进一步明确组合条件,即限制生成能力,使“黑夜怕人”这样的句子不出现。从另一个角度看,Harris(1946)的语类生成模型和 Chomsky(1957)的短语规则生成能力又太弱,比如 Chomsky(1957)的短语改写规则中,VP 只有下面的改写规则:

VP → Verb + NP

这一规则不能生成下面的句子:

张三送李四一块手表 (NP1 + V + NP2 + NP3)

张三知道巴黎在法国 (NP1 + V + S)

张三喜欢学英语 (NP1 + V + VP)

.....

实际上,动词 VP 和它周围的直接成分关系相当复杂,所以必须分成很多小类。Chomsky(1965)采取的办法是对动词进行次范畴化。

Chomsky(1965, 2. 3. 4, p94, 规则 40)讨论了严格次范畴化规则(strict subcategorization rules),这些规则涉及用其他范畴框架来给动词次范畴化,把动词分析成复合符号(complex symbol),于是动词就分成了很多小类,不同小类的动词有不同的搭配语境,这些信息需要包括在 lexicon 里(Chomsky, 2. 3. 4, 规则 41, p94)。比如:

eat: [+v, +__ NP] (eat 可以跟一个名词短语 NP)

believe: [+V, +__ NP/ +__ that S'] (believe 可以跟一个名词短语 NP, 或一个句子 S)

动词的严格次范畴化具有普遍的意义,显示了一个动词可以和哪些语类搭配。

具有相同搭配框架的动词,共现情况并不相同,以 NP1 + V + NP2 为例:

大人吓唬小孩

老师吓唬学生

· 空气吓唬小孩

· 老师吓唬汽水

以上例子显示“吓唬”前后的 NP 是有选择限制的。一旦前后的

NP 规定下来,V 也是有选择限制的:

大人吓唬小孩

大人骂小孩

* 大人写小孩

* 大人节约小孩

面对这样的问题,我们又回到了前一节所讨论的动词的选择规则问题。Chomsky(1965,2.3.4,p95,规则 42)在动词严格次范畴的前提下讨论了选择规则(selectional rules),对于出现在相同框架中的动词,选择规则进一步把这些动词分析成复合符号,复合符号包含一组句法特征,只有满足这些句法特征的词才可以进入相应的搭配框架。

由于动词的次范畴化涉及名词或其他词所组成的上下文,所以被 Chomsky 称为“受上下文制约的次范畴化规则”(context-sensitive subcategorization rules)。一般地说,次范畴信息实际上是上下文信息(contextual information),因为它规定了一个词出现的上下文。

概括地说,Chomsky 的次范畴化有以下一些特点:

1. 次范畴化及其语义特征的分析(p82)是以解释组合规则为目的的。这就促使人们更多地要从组合关系角度去给词分小类。比如在“吃饭”这样的组合中,给“饭”分出“饮食”的特征,对组合是有意义的,但分出“主食”的特征,对组合的意义不大。

2. Chomsky 先对名词次范畴化,再对动词次范畴化(p86,p90)。具体地说,先描写名词的句法特征(相当于语义特征),就可以比较好地确定动词的搭配框架。这就是把动词分成了不同

的小类,每一个小类都有自己的搭配框架,由此达到了限制生成能力过强的目的。这正是次范畴化(subcategorization)的含义。从这种意义上说,名词次范畴化不考虑上下文,动词次范畴化要考虑上下文。

3. 句法中心论。标准理论的深层结构是句法深层结构。NP、VP 等都是深层结构中的根本范畴。语义是解释性的(p75)。

Chomsky(1965, 2. 3. 4, p101)进一步讨论了介词短语(Prepositional-Phrase)对动词的次范畴化作用。Chomsky(1965, 2. 3. 4, p101 — 103)区分了 Prepositional-Phrases 和 Place and Time adverbials,前者作用于具体动词,涉及动词的次范畴化,后者跟整个动词短语或句子联系,对动词的次范畴化没有影响。Chomsky(1965, 2. 3. 4, p102)说:

It will follow, then, that Verbs will be subcategorized with respect to Verbal Complements, but not with respect to Verb Phrase Complements.

7.2.3 次范畴化和平行周遍条件

Chomsky(1965)把次范畴化和 branching 区别开。Chomsky(1965, 2. 3. 2, p79)明确指出下面的区别:

Note that this information concerns subcategorization rather than“branching”(that is, analysis of a category into a sequence of categories, as when S is analyzed into NP+Aux+VP, or NP into Det+N).

关于 branching rules 和 subcategorization rules 的区别,还可以

进一步参考以下说明(Chomsky 1965, 4.1, p112):

Thus a branching rule analyzes a category symbol A into a string of (one or more) symbols each of which is either a terminal symbol or a nonterminal category symbol. A subcategorization rule, on the other hand, introduces syntactic features, and thus forms or extends a complex symbol.

其实这两个问题是同一个问题的不同方面。branching rule 本质上就是短语结构规则,是在已经有的语类基础上组合的类改写成直接成分的类。Chomsky(1957)的短语结构规则或者说 Harris(1946)的语类规则之所以不充分,是因为早期的语类工作分类不细。次范畴化实际上是在考虑词的搭配框架,和上下文语境的语义特征,实际上是把词分成更小的类。所以次范畴化本质上是给原来的语类作了更细的分类,在这一基础上再建立语类规则,这个时候 branching 的组合条件就更严格了,对线性组合规则的描写就更加充分。Chomsky(1964)认为结构语言学是分类语法,实际上标准理论中的基础部分的范畴和次范畴的工作都是在做分类的工作。可以看出,分类语法分析实际上是语法分析必不可少的一个层面。当然,新的更细致的分类需要以解释组合规则为目的,而不是盲目的分小类。

无论是对名词展开次范畴化,还是对动词展开次范畴化,从根本上可以概括成本章开始提到的两个问题,一个是组合框架的问题,一个是组合框架中语类的选择问题。

对语类次范畴化以后,句子的生成能力受到语义条件限制,生成能力虽然减弱了,但生成符合语法的句子能力增强了。下面具体分析语类次范畴化以后句子生成的机制:

短语规则：

$$S \rightarrow NP + VP^{\textcircled{1}}$$

$$VP \rightarrow V + NP/PP/S'$$

词库(lexicon)：

eat: [+v, + ___ NP(+可食用的)]

think: [+V, + ___ that S']

frighten: [+V, + ___ NP(+人类的),]

.....

sincerity[+N, -可数的, +抽象的]

boy[+N, +可数的, +普通的, +动物的, +人类的]

.....

从优先选择短语规则的角度看,如果在短语规则中选择了

$$VP \rightarrow V + NP$$

就只能在词典中选择 eat, frighten 等词项而不能选择 think 等词项。如果选择的 frighten 词项,宾语就只能选择具有“+人类的”这一特征的词,即只能选择 boy 等词而不能选择 sincerity 等词。这样就能生成:

sincerity may frighten the boy.

① 为简单起见,这里不考虑助动词。但是必须注意,标准理论为了给自然语言提供严格的形式化描写,任何一细节都是不能漏掉的。请参考原文。

在后面的讨论中我们会看到,句子的生成过程存在句法优先和词库优先的区别。上面是句法优先的思路。我们也可以采取词库优先的思路。词库优先是先选词,如果在基础部分的词典中选择了 frighten,下一步就只能选择规则:

$$VP \rightarrow V + NP$$

往下的步骤和句法优先是一样的。无论是句法优先还是词汇优先,都需要做名词和动词的次范畴化的工作。

Chomsky(1965)明确把句法的基础部分分成改写规则部分和词库部分。Chomsky(1965, 2.4.3, p120)认为次范畴规则可以置于词库部分:

(One might propose, alternatively, that the subcategorization rules be eliminated from the system of rewriting rules entirely and be assigned, in effect, to the lexicon. In fact, this is a perfectly feasible suggestion. (1965, 2.4.3, p120))

这种思路的实质是:

1. 动词词库中标注好搭配框架和框架中词的次范畴特征。
2. 名词词库中必须有词的次范畴特征。

根据词汇插入规则,如果名词的次范畴特征和动词搭配项的次范畴特征不矛盾,就可以插入动词前后的名词位置,生成合法的句子。我们认为这不是处理好坏的问题,动词的搭配框架或次范畴特征是个别事实,是动词的个性,都应该存放在词库中。Chomsky 的处理方法是有道理的。后来的词汇功能语法实际上是这一思路的扩展。引入语义,展开次范畴化分析,目标

是要解决生成能力太强的问题,这种研究在理论上的一个重要结果是,很多语言知识是在词库中,不是在规则中。很多生成能力太强的规则,那是因为把次范畴描写的内容放到了规则上。从根本上说,不规则现象都要放在词库中。

次范畴化理论区分了规则搭配和不规则搭配,具有重要的理论意义,但哪些属于规则搭配,哪些属于不规则搭配,Chomsky 没有提出严格的判定标准,所以哪些应该存放在词库中,哪些应该存放在规则库中,当时并不清楚。比如,Chomsky(1965, 2.3.4, 规则 52)进一步把 Adverbial 等制约动词次范畴化的条件放入规则,并进一步区分了动词补语(Verbal Complements)和动词短语补语(Verb Phrase Complements),前者会引起 Verbs 的次范畴化,后者不会。观察 Verbal Complements 的例子(Chomsky, 2.3.4, p103):

he argued *with John* (*about politics*)

he aimed (the gun) *at John*

he talked *about Greece*

he ran *after John*

he decided *on a new course of action*

Verb Phrase Complements 的例子:

John died *in England*

John played *Othello in England*

John always runs *after dinner*

按照 Chomsky 的解释,动词补语是动词的补语,动词短语补语是动词短语的补语或句子的补语,这种区分是后来 XP 语杠理论的一个基础。但是如何判定一个补语是动词补语还是动词短

语补语, Chomsky 没有给出标准, 所以他对次范畴化的处理标准还不明确。Chomsky(1965, 4.2, p165)假定:

Furthermore, let us make the empirical assumption that a grammar is more highly valued if the lexical entries contain few positively specified strict subcategorization features and many positively specified selectional features.

这是要求尽可能把严格次范畴化的信息放在句法中, 但并没有提供一个严格的标准。我们认为问题可以通过平行周遍条件来解决。如果 $X + \text{Comp}$ 的搭配是平行周遍的, 补语 Comp 就是动词短语的补语。例如:

John died in England
John lived in England
John worked in England
John studied in England
John called in England
John lost in England
John visited in England
.....

凡是有处所特征的动词 V, 只要语境允许, 都可以和 in England 组合。而处所特征不需要记忆, 可以从认知上得到判定。对于动词补语(Verbal Complements), 通常找不出这样的平行周遍条件, 动词和补语的组合是高度个性化的。

我们说动词短语补语(Verb Phrase Complements)前面的动词可以在某个平行条件下周遍一类动词, 即满足这样一种关系:

V + Prep. + NP
 V1 + Prep. + NP
 V2 + Prep. + NP
 V3 + Prep. + NP

 Vn + Prep. + NP

动词补语(Verbal Complements)前面的动词尽管不能进行平行周遍替换,但仍然是规则组合而不是固定搭配,因为其中的 NP 是可以平行周遍替换的,符合本书前面 3.6 节给出的多个语素组合的规则条件。比如:

[V][Prep. + NP]
 He argued with the students
 He argued with the clever students
 He argued with the clever students of Peking University

由于 with N 不能进行平行周遍替换, argue with N 需要进入词库。一般地说,对于格式:

V + Comp 或 [V] + [Prep. + NP]

如果是动词短语和补语的关系,可平行周遍替换的是其中的直接成分 V,如果是动词和补足语的关系,可以平行周遍替换的是其中的直接成分 Comp 所统辖(dominate)的更小成分 NP。正是因为动词补语和动词的这种非平行周遍的替换关系,所以这种组合格式需要进入词库。

考虑“打”的一个义项：

打水

“打”在这里有“取、买”的含义。词库中首先需要的次范畴信息是：

打:[NP1 v NP2]

NP1 的次范畴特征必须包含“人”的特征,如果允许拟人化,应该包含“动物”特征。从下面的平行实例看,NP2 应该包含“液体、散装”的特征:

打水、打酱油、打牛奶、打醋、打汽油、打酒……

· 打“娃哈哈”、打“长城干红”

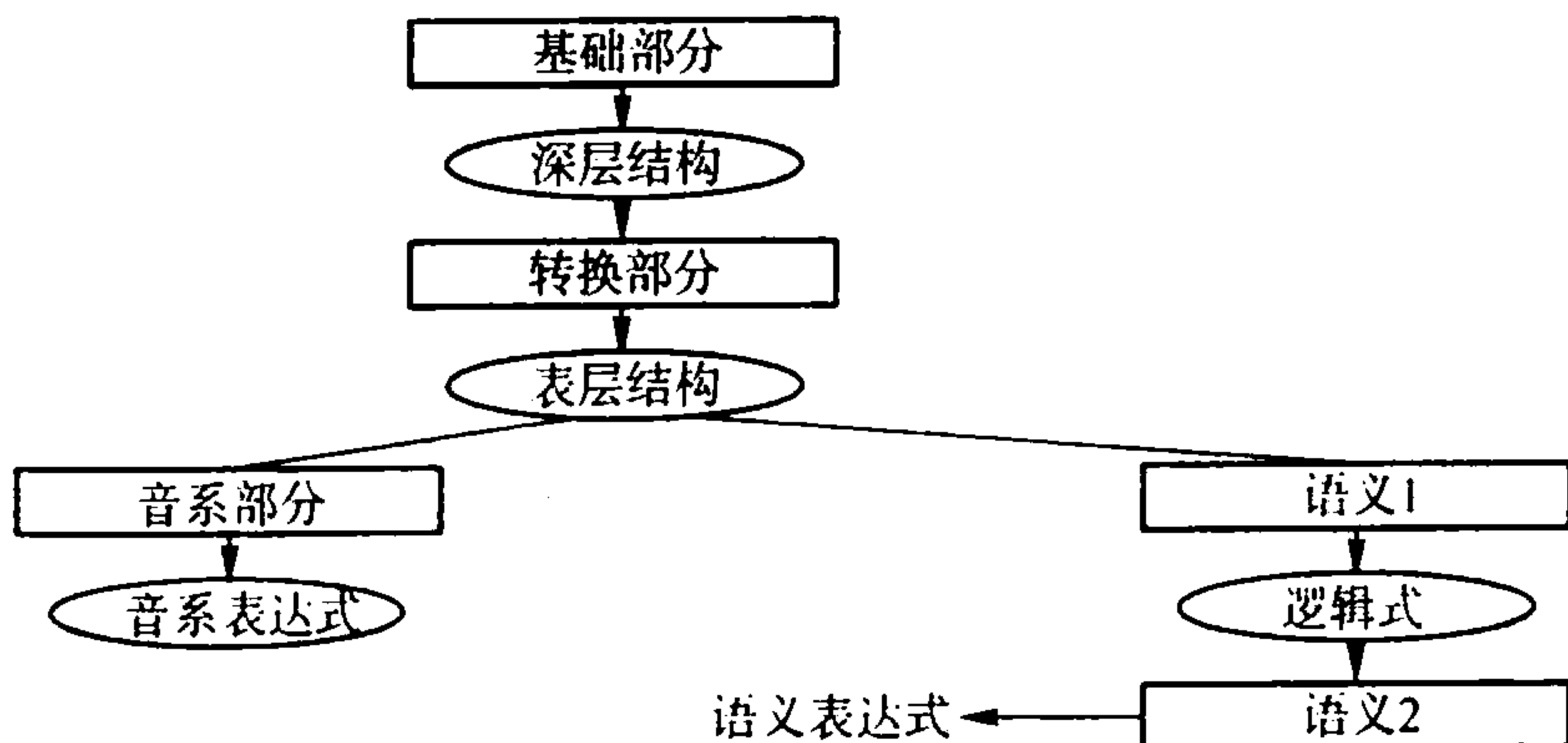
如果这种平行条件是周遍的,这些组合就不需要存放在词库中,词库中有“NP1 打 NP2”就够了。

7.3 深层结构及其线性特征

词的组合框架以及词的选择限制属于语义问题还是句法问题,不同的学者有不同的看法,Chomsky(1965)以次范畴化规则和选择规则来解决这类问题,基本上把这类问题看成是句法问题,Chomsky(1965,2.3.1,p75)也提到这类问题是属于语义还是句法并不是很清楚。不过 Chomsky 对这类问题的解释集中在深层结构上,或基础部分,不是在转换规则上,这和很多学者是一致的。现在要讨论的问题是,深层结构是否存在,是否是一个独立的部分。承认深层结构就意味着区分深层结构和表层结

构。Chomsky(1965)以前已经有人提出了深层结构和表层结构的区分问题,比如 Wittgenstein(1953, p168)的 depth grammar and surface grammar、Hockett(1958)的 depth grammar and surface grammar、Postal(1964)的 underlying structure and superficial structure。Chomsky(1965)的转换生成语法标准理论系统讨论了深层结构(deep structure)、表层结构(surface structure)以及两者的关系。深层结构是由基础部分生成出来的,在标准理论中有重要价值。深层结构和表层结构的区别是标准理论中的重要区别。此后 Chomsky 一直坚持认为有一个和表层结构相对的深层结构的层面,只是对深层结构的认识有一些改变。

标准理论的深层结构有一个重要作用就是解释语义。Chomsky(1972)在扩展的标准理论(Extended Standard Theory, EST)中,认为跟量词、否定等相关的意义可以在表层结构中解释,深层结构解释的是语法关系的意义。扩展标准理论(EST)要在深层结构和表层结构两部分解释语义,理论上不太统一,描写也比较复杂。1977年,Chomsky 又对扩展的标准理论进行了修正,代表作是《关于形式和解释的论文集》(*Essays on Form and Interpretation*, 1977)。这一次修正被称为“修正的扩展的标准理论”(Revised Extended Standard Theory, 简称 REST)。这一次修正的特点是把语义解释全部放到表层结构。结构如下:



从结构框架可以看出,REST 有两个特点:

- 1. 在表层结构解释语义,给出语义解释的逻辑表达式。
- 2. 语法研究只处理语义 1,即和逻辑式有关的语义关系,包括语义格以及生成语义学中不分解词义的语义表达式。

为了在表层结构解释语义,语迹的概念出现了。比如,对于下面的序列:

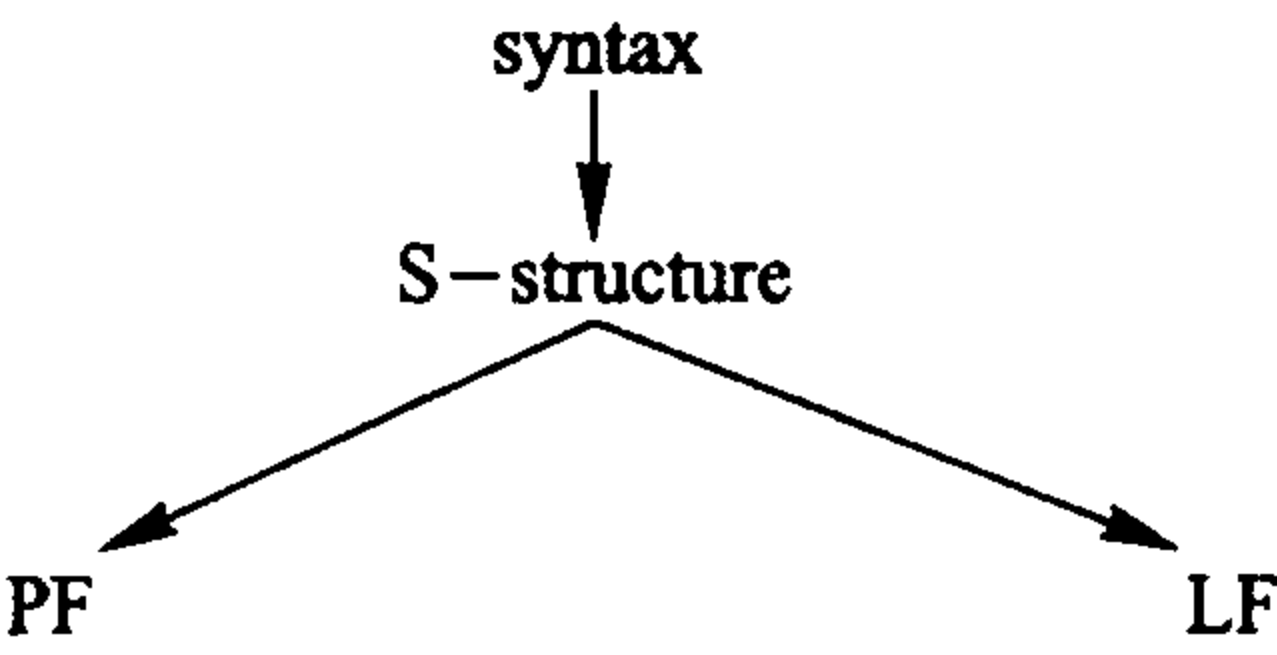
Who will you visit

可以看成是下面移位并留下语迹的结果:

You will visit him
Will_j you t_j visit him
Who_i will_j you t_j visit t_i

通过几次修正,语法规则系统基本上确定下来了。在 Chomsky (1981)支配和约束理论中,意义基本上在表层结构解释,深层结构仍然存在,是通过转换映射到表层结构的基础:

three fundamental of UG



举例来说,对于下面的一个实例(Chomsky,1981,p33):

it is unclear who to see

各个层面的相互关系是：

1. it is unclear [_{s1} who [_s to see]] (表层结构, 和可观察的句子基本是一致的。)
2. it is unclear [_{s1} who_i [_s PRO to see t_i]] (S-结构, 有语迹)
3. it is unclear [_{s1} COMP [_s PRO to see who]] (D-结构, 尚未经过转换, 所以不存在语迹)
4. it is unclear [_{s1} for which person x [_s PRO to see x]] (LF, 即逻辑形式, 是解释性的, 由 S-结构派生出来的。)

以上的 D-结构(相当于深层结构, 有微小的区别)中并不解释语义, 所以看不出 see 的施事和受事指称关系。S-结构相当于表层结构, 但已经标注了语迹, 所以能从下标看出 who_i 和 t_i 指称相同。LF 则根据 S-结构进一步解释施事 PRO, 受事是 which person x。

Chomsky(1995)的最简方案(the minimalist program)相对 Chomsky(1981)的观念有一些变化。最简方案取消了 D-结构, 词库实际上取代了深层结构, 句子的各种差异都直接从词库的投影中形成。句子通过直接在词库中选词并通过合并(merge)和移动产生大显型(Spell out), 然后在 PF 部分借助发音感知系统得到语音解释, 在 LF 部分借助概念系统得到语义解释。显然, Chomsky(1995)的最简方案吸收了词汇功能语法的从词汇入手的核心思路。

最简方案还是高度探索性的, Chomsky(1995)在最简方案中没有提深层结构, 也没有给出取消深层结构的论证理由。我们后面要论证, 由于最简方案仍然存在一个转换层面, 转换赖以操作的线性结构实际上相当于深层结构。

深层结构在很多研究中仍然有重要地位。下面我们来分析深层结构的性质。

7.3.1 深层结构的存在依据

深层结构不完全是一种可观察的结构,如果能够不用深层结构,当然对理论的实证性有好处。目前很多关于深层结构存在理由的分析,并不是很充分的,所提到的深层结构好些都可以通过可观察结构来解释。比如 Chomsky(1965, p22)在论证描写充分性时所举的例子:

(6) I persuaded John to leave

(7) I expected John to leave

Chomsky 指出以前还没有语法指出这里的区别:

The first impression of the hearer may be that these sentences receive the same structural analysis. Even fairly careful thought may fail to show him that his internalized grammar assigns very different syntactic descriptions to these sentences. In fact, so far as I have been able to discover, no English grammar has pointed out the fundamental distinction between these two constructions.

Chomsky(1965, p22—23)通过以下分析证明了上面(6)和(7)深层结构不同:

(8)(i) I persuaded a specialist to examine John

(ii) I persuaded John to be examined by a specialist

(9)(i) I expected a specialist to examine John

(ii) I expected John to be examined by a specialist

这 4 个句子的深层结构分别是：

- (10) (i) Noun Phrase-Verb-Noun Phrase-Sentence
 (I-persuaded-a specialist-a specialist will examine John)
 (ii) Noun Phrase-Verb-Noun Phrase-Sentence
 (I-persuaded John-a specialist will examine John)
- (11) (i) Noun Phrase-Verb-Sentence
 (I-expected-a specialist will examine John)
 (ii) Noun Phrase-Verb-Sentence
 (I-expected-a specialist will examine John)

Chomsky 这里的观察和描写是相当细致的, 不过, 既然 Chomsky(1965) 已经给出了次范畴规则, 以上的歧义实际上可以通过次范畴规则来解决, 所以 Chomsky 在这里并没有给出深层结构存在的理由。

转换也不是深层结构存在的理由, 因为句子的转换也可以在可观察的句子上进行。Harris(1952、1957、1965) 采取的就是这种办法。朱德熙(1986) 在汉语中的变换分析也是这样的方法。Chomsky(1957) 把一些使用转换规则很少的句子称为核心句, 其他句子可以在核心的基础上转换出来, 也是这种思路。

前面讨论过, 通过短语结构规则不能生成所有的句子, 有些句子的生成, 必须要有转换规则和音位规则对短语结构规则的句子进行处理。在汉语中, 下面的句子“我把球扔了”通常认为要通过转换:

我扔了球

→我把球扔了(把字转换)

→我把球扔了(变调, 上声变调, 轻声变调)

因此我们可以采取一种描写方法,说“我把球扔了”是从“我扔了球”转换过来的,不过转换前和转换后的这两个句子都是可以观察到的事实。

Chomsky 和 Lasnik(1977)谈到语迹会阻止语音的融合。考虑下面的实例:

I want to see him

→ Who_i do you want to see t_j (移位,产生语迹)

→ who do you wanna see (产生语音合并 want to→wanna)

I want Jack to see John

→ Who_i do you want t_j to see John (移位,产生语迹)

→ * Who_i do you wanna see John (由于有语迹存在,无法完成语音合并 want t_j to→wanna)

We should have come

→ Should_i we t_j have come (移位,产生语迹)

→ * Should_i we've come? (语迹阻止 we have→we've 的融合过程)

语迹存在的证据似乎可以进一步论证深层结构的存在。不过语迹是可以还原的,语迹还原以后的组合通常是可以观察到的组合,因此严格地说语迹的存在只能证明确实产生了移动现象,但并不必然证明深层结构的存在。

如果自然语言中所有的句子在转换前和转换后都可以观察到,就可以不必假设有深层结构。但是,对于下面的句子:

将李四的军

这个句子不可能通过连接和替换得到,只能通过转换得到。在转换之前,并不存在一个“将军”和“李四”直接相组合的可观察句,这就需要承认一个深层结构:

将军 + 李四

我们把转换规则和音系规则统一称为变形规则。变形规则没有递归性,是对短语结构规则生成的语符序列的处理。既然任何语言都必须要有变形规则,就必须有变形前的结构,如果使用变形规则以前的结构包含有不可观察的,就说明语言中有抽象的形式或底层形式,这就是深层结构存在的本质。

为什么不直接从词库中的结构生成“将李四的军”呢?其实这也是一种可以考虑的方案,我们可以在词库中注明“将军”后面不能直接跟受事 NP,必须有“将 NP 的军”这样一种映射方式把词库中将军的受事映射到表层结构。更一般地说,对于 ab 两个单音节语素构成的动词 V,可以有两种和受事 NP 组合的方式:

$VP \rightarrow V + NP$

$VP \rightarrow a + NP + \text{的} + b$

然后在语符库中注明“将军”一类动词属于第二种类型。这种解释方式统一性不够强,而且这种词库标注只能解决这个实例本身,有些不可观察结构不能够用词库标注的方式来解决。前面我们在论证转换规则的必要性时给出过“我知道他去过上海了”,并分析过这个句子的层次可能有三种情况:

我(知道[他去过上海了]) (他去过上海了,我知道)

我([知道他去过上海]了) (他去过上海,我知道了)

我(知道[他去过上海了]了) (他去过上海了,我知道了)

第三种情况显然是不可观察的,而且也不可能在词库中标注。这类现象是深层结构存在的关键证据。

Chomsky(1995)在最简单方案中取消了深层结构。动词及其语义格信息全部存放在词库中,句子的线性组合直接通过词与词的合并(merge)从词库中派生。但要彻底取消深层结构是有一定困难的,因为自然语言中存在“我(知道[他去过上海了]了)”这类组合,合并(merge)不可能生成这样的不可观察线性序列。从递归性的角度看,通过合并生成线性序列也有问题。这需要再来回顾一下深层结构的本质。我们前面提到,深层结构是线性结构的一种,即一种不可观察的线性结构。因此,像可观察线性结构一样,深层结构本质上也是在连接和替换的基础上形成的结构,具体地说,先通过语类的连接和替换生成有层次的语类序列,再通过词汇的插入生成深层结构。因此,深层结构的本质是单位有层次的线性组合。Chomsky 最简方案取消深层结构以后,句子的生成过程是先在词库中取词进行合并,比如选择“写”和“诗”进行合并有:

写诗

前面我们分析过,初始连接的线性顺序是相当确定的,一般不会出现“诗写”的顺序,因此,即使我们取消了深层结构这个术语,初始连接的相对稳定的线性顺序并没有改变,所以深层结构的实质并没有改变。

更重要的是,我们前面讨论过,深层结构本质上就是通过初始连接、替换形成的有语类标记的线性结构,并且只有这种结构才有递归性。考虑词库中“砸”的语义结构框架:

$[NP^{施事}] + 砸 + [NP^{受事}]$

如果要扩展替换 NP, 必须要有这样的规则:

$NP \rightarrow Spec + N^1$

$N^1 \rightarrow Subj + N^1$

$N^1 \rightarrow Comp + N$

才能生成这样的句子:

小邵砸了两个白色的玻璃杯子

这一套改写规则包含递归规则 $N^1 \rightarrow Subj + N^1$, 因此名词短语的改写规则不可能在词典中解决, 可见名词短语的改写规则是肯定需要的。更一般地说, 语类规则是必要的。这里的论证过程是以扩展替换为前提的。最简方案线性结构的生成采取的是合并(merge), 是否就可以绕过短语改写规则呢? 根据我们前面讨论过的语类规约原则, 合并也不能绕过短语改写规则。比如, 上面以“杯子”为中心的名词短语要能够扩展替换, 必须要满足语类规约原则, 即扩展出来的短语必须是 NP, 这就决定了杯子前面的附加语必须受到限制:

小邵砸了($_{NP}$ [$_N$ 玻璃][$_N$ 杯子])

小邵砸了($_{NP}$ [$_A$ 小][$_N$ 杯子])

*小邵砸了($_{VP}$ [$_V$ 买][$_N$ 杯子])

以上实例证明当其他成分和“杯子”进行合并的时候, 我们必须知道合并前和合并后的语类关系:

$$N+N \rightarrow NP$$

$$A+N \rightarrow NP$$

$$V+N \rightarrow VP \text{ (不包括“烤白薯”类)}$$

而这种语类推导正是改写规则的逆过程：

$$NP \rightarrow N+N$$

$$NP \rightarrow A+N$$

$$VP \rightarrow V+N$$

总之,无论是替换扩展还是 Chomsky 最简方案中的合并扩展,都不能绕过语类规约原则,而改写规则正是语类规约原则的体现,因此合并也不可能绕过改写规则。这就决定了合并生成的结构仍然是一种符合一定语类排列的线性结构。合并过程也要受到语类限制,得到的同样是语类线性结构。由于语类规约原则是递归原则,不可能放在词库中,所以合并把线性组合的全部信息存放在词库中是不可能的。

当然,Chomsky(1995)的合并是指具体的词和具体的词相组合,有些类似我们前面讨论的线性程序的叠加。如果严格遵循具体词的合并,是否可以绕过语类规约原则? Chomsky 没有论证过。我们认为不能。合并如果要绕过语类规约,就不得不采取记忆或词典方式记录组合的类。尽管一个语言的单位是有限的,因此两个单位的组合也是有限的。比如,如果有 10000 个单位,两个单位的组合不会超过 10000^2 。从理论上说,可以记住两个单位的组合的语类,比如:

n: 大杯子、小杯子、玻璃杯子、塑料杯子……

a: 很好、非常好、相当好、格外好……

v: 经常来、明天来、一起来、一定来……

但是,组合的长度不只是两个。三个单位的组合的类是否也要记住? n 个单位的组合的类是否也要记住? 这是不可能的,必须要依赖语类规约原则。比如下面的 n 个名词的组合:

N	N	N	N	N...	N.
中国	上海	农业	科学	技术	学院

最终是通过 $N+N \rightarrow N$ 的方式规约为 N ,并不是通过记忆记住的。下面两个组合只有一个成立:

吃_[NP 烧土豆]
 *吃_[VP 卖土豆]

那是因为“烧土豆”可以规约成 NP,“买土豆”不能。

再来看合并的层次问题。对于下面的实例:

三个研究所的所长

这里的歧义是因为层次的关系,这是学界都熟悉的。要区别这里的歧义,采取合并的方式就必须有顺序进行合并:

([三个研究所]的)所长
 三个([研究所的]所长)

因此,合并也不能绕开层次问题。

把上面的讨论归纳起来,合并不能绕开语类规约,也不能绕开层次。现在再来看早期深层结构的实质。短语结构规则的改写结果生成短语标记(phrase marker),插入词汇就构成深层结构。这时的深层结构有下面两个特征:

1. 直接成分的线性组合(层次)
2. 每个成分的语类标记
3. 具体的单位

由于最简方案中的合并程序并没有绕过以上特征,所以合并所生成的结果本质上具有上述特征,因此跟深层结构的含义是相同的。

如果承认词所组成的线性结构受到语类和层次的限制,也就承认了早期的短语标记(phrase marker),也就承认了深层结构。当然我们可以不用深层结构这个术语,而用线性结构这个术语,但并没有改变问题的实质,因为由线性程序生成的线性结构有两种,一种是直接可观察的组合,比如:

拿一本书

还有一种是不可观察的组合,需要转换才能成为可观察形式:

我(知道[他去过上海了])了)→我知道他去过上海了

给出一个深层结构层面,有利于解释第2类句子。当然,我们仍然可以不用深层结构这一概念,而直接用线性结构这一概念。但需要承认,有些线性结构具有抽象性或不可观察的特点。

在音系层面,相当于深层结构的底层形式也是存在的。前面曾经讨论过的北京话儿化就属于这类情况。比如“盖儿”的派生过程:

$$\text{kan}^{55} + \text{er}^{35} \rightarrow \text{kar}^{55}$$

实际语料中 $\text{kan}^{55} - \text{er}^{35}$ 的可观察形式是找不到的。北京话连上

变调也属于这样的情况,一般人都承认下面的两个组合并不一样:

土改 $214+214\rightarrow35-214$

涂改 $35+214\rightarrow35-214$

这里的不一样就是因为底层组合不一样,表层是一样的。这里的底层组合 $214+214$ 是不可观察调型。

7.3.2 深层结构的确定

既然不得不承认自然语言中有不可观察组合,就需要承认深层结构,需要有一种统一的生成程序来确定深层结构。我们认为线性程序就是这样的生成程序,而深层结构就是可以用线性程序解释和生成的结构。前面提到,从线性生成程序的角度看,也有两种线性结构,一种是在初始连接的基础上通过扩展替换得到的,我们可以把这种线性结构作为深层结构,这时转换规则会增加。用这种线性结构做深层结构,题元的位置比较固定,因为通过替换扩展出来的线性序列总是可以还原到初始组合,而在初始组合中,语法结构成分、语义结构成分是相对固定的,这就为在深层结构的位置上确定题元提供了基础。比如:

我们已经读过《阿 Q 正传》了

通过替换还原,我们可以得到:

我们读《阿 Q 正传》

这是初始组合。由于用替换结构能够为转换提供一个参照结构,题元的位置在替换结构中是相对固定的,所以用替换结构作

深层结构或线性结构有价值。尤其是对“出太阳”这样的片断，“太阳”是施事还是受事不容易确定，用替换结构作深层结构有着更进一步的意义。

另一种线性结构是允许使用叠加规则形成的。比如：

《阿 Q 正传》我们已经读过了

其线性连接过程是：

读

→读[过]

→读过[了]

→[已经]读过了

→[我们]已经读过了

→[《阿 Q 正传》]我们已经读过了

方括号中的成分是层层叠加上来的成分。在第 2 章已经讨论过，通过这种叠加生成的线性序列也有层次，但不能还原成初始组合，所以论元位置相对来说是不固定的。

叠加是比替换更强的生成方式，只要是替换得到的序列，用叠加一定能够得到，但用叠加得到的序列，替换不一定能够得到，比如“一匹马”就不是替换能够得到的。用扩展替换得到序列可以通过简缩替换得到原形，通过叠加得到序列不一定能够通过简缩替换得到原形。因此选择替换结构还是叠加结构作深层结构或线性结构，转换规则的多少会有所不同。如果在替换结构的基础上使用转换，只有叠加才能形成的线性结构都作为转换结构，转换规则就会增加。如果在叠加的基础上使用转换规则，转换规则就会减少。从题元解释的角度看，用替换结构作深层结构更为优化。另外，用叠加来限定线性结构或深层结构，

有时候不好判定层次。比如：

杯子/汪锋砸了

如果认为是叠加的结果,就需要承认[汪锋砸了]是独立片断,这就相当勉强。下面两个实例确实是平行的:

杯子/汪锋砸了

《阿 Q 正传》/我们读过

但如果认为第 1 个句子是转换来的,第 2 个句子是叠加来的,就不好解释这里的平行关系。从这个角度看,用替换限制线性结构或深层结构也是有道理的。

7.3.3 深层结构的线性顺序

Chomsky(1957)的短语结构规则已经假定了语类组合的线性顺序:

$$\begin{aligned} (69) \quad S &\rightarrow NP + VP \\ VP &\rightarrow V + NP \end{aligned}$$

即句子的主语在 VP 的前面,宾语在 V 的后面。另一种是不给出线性次序:

$$\begin{aligned} (70) \quad S &\rightarrow \{NP, VP\} \\ VP &\rightarrow \{V, NP\} \end{aligned}$$

像(69)这样的系统称为 contatenation-systems(连续系统),像(70)这样的系统称为 set-systems(集合系统)。连续系统就是

线性系统。Chomsky(1965, 2.4.4, p124)讨论了 NP 和 VP 的顺序问题,认为超越线性或连续系统的语法关系研究是不可能的,这一态度代表了一种重要的理论取向。我们在把原文引证如下:

Proponents of *set-systems* such as (70) have argued that such systems are more “abstract” than *concatenation-systems* such as (69), and can lead to a study of grammatical relation that is independent of order, this being a phenomenon that belongs only to surface structure. The greater abstractness of set-systems, so far as grammatical relations are concerned, is a myth.

Chomsky(1965, p125)的基本理由是:

Hence, it seems that a set-system such as (70) must be supplemented by two sets of rules. The first set will assign an intrinsic order to the elements of the underlying unordered Phrase-marker (that is, it will label the lines of the tree-diagrams representing these structures). The second set of rules will be grammatical transformations applying in sequence to generate surface structures in the familiar way. The first set of rules simply converts a set-system into a concatenation-system. It provides the base Phrase-markers required for the application of the sequences of transformations that ultimately form surface structures. There is no evidence at all to suggest that either of these steps can be omitted in the case of natural languages.

归纳起来,Chomsky 认为范畴部分的规则起两个作用:

1. 确定了语法关系。
2. 确定了深层结构中要素的顺序。

我们可以把 Chomsky 的深层结构看成线性深层结构的一种,这种线性深层结构并不给和动词直接发生结构关系的名词或谓词定性,即不用“施事、受事”等定性概念,这和后来 Fillmore 的格不完全相同,代表了两种思路。我们把 Fillmore 的格语法看成是非线性深层结构的一种。沈阳(1994)在汉语研究中的方法和 Chomsky(1965)的思路是一致的。

不过 Chomsky 的论述还不足以证明深层结构是连续系统或线性系统。后来 Chomsky 和 Lasnik(1977)在研究英语的时候发现,语迹前后语音有不产生并合的现象。我们在本书前面 7.3.1 节中已经分析过类似的例子,这一事实也只能证明转换前的结构是线性的,不能证明深层结构的线性关系一定存在。

把深层结构作为连续系统的根本理由应该是我们前面提到的初始连接的性质。考虑下面两个语符的初始连接。

学英语、走路

* 英语学、* 路走

初始连接都是 V+N 顺序。再看 NP+VP 的初始连接:

学生学英语、孩子走路

* 英语学生学、* 路孩子走

* 学生英语学、* 孩子路走

也都是有顺序的。一般地说,初始连接以及在初始连接基础上

通过替换形成的结构都是有顺序的。有些话语的生成过程必须假定有不可观察的初始连接,然后再使用变形规则得到。前面我们说“我知道他去过上海了”有一种分析是“我(知道[他去过上海了]了)”,就属于这样的情况。

当然,也可以把无线性顺序的集合系统看成是比深层结构更初始的语义结构系统,这时就可以完全不考虑顺序。从语法优化角度看,线性深层结构有利于判定动词前后题元的性质,或者说线性深层结构已经在很大程度上规定了题元的性质。题元移动时通常都有表层的句法标记,所以仍然能够识别题元。集合系统还必须假定我们目前对论元的性质已经有了一个统一的分类标准,但这项工作目前远远没有完成。“出太阳”中的“太阳”是施事、受事还是客事,目前还没有统一的意见。但在初始连接中应该是“出太阳”而不是“太阳出”,这是可以肯定的。

思考问题

1. 深层结构的存在依据是什么?
2. 动词的次范畴化是否需要考虑平行周遍条件?

8. 语义组合关系的非线性化及其功能分析

Chomsky(1965)用语类所处的位置来解释语法功能层面的信息,句子 S 的直接成分 NP 就是主语,谓语 VP 的直接成分 NP 就是宾语,这样主语、宾语等概念就可以通过语类推导出来。Chomsky(1965)又用动词的次范畴化来描述语义功能层面的内容,基本上不提语义结构关系、语义格或语义功能等概念,线性顺序在 Chomsky(1965)的动词次范畴描述中起了关键作用。拿汉语来说,“吃”的严格次范畴框架是:

吃: $NP_1_V_NP_2$

即“吃”是出现在两个名词之间的动词,其中 NP1 应该是动物, NP2 应该是“可入口的”。从语义结构关系的角度看,这里的 NP1 就是施事, NP2 就是受事。从 Chomsky(1965)的行文来分析,Chomsky 没有提语义格的概念,目的之一是为了追求形式化,限制用过多的初始概念,尤其是语义功能概念。

Fillmore(1968)的格语法(Case Grammar)以一种不同于 Chomsky(1965)的方式来分析语义关系。Fillmore 首先引入了一些反映语义角色的初始概念,如“施事、受事、工具、与事”等,格语法中动词和周围作语义格的名词的关系并没有顺序,只是词汇插入后才有了顺序,或者说是在语义结构实现为表层结构的过程中才有了顺序。格语法实际上是在超越深层结构的线性特征。为句子的语义解释找出一个更深的,非线性的底层结构。这实际上就是 Chomsky(1965)提到的集合系统

(set-system)。

这里我们可以引出一个原则问题,引入语义结构以后,语义格的组合是否遵循线性规则?

8.1 格语法的生成机制

主—谓、述—宾这样一些语法结构关系和施事—行为、行为—受事这样一些语义结构关系,是结构语言学以前或前结构语言学关心的问题,这是一种功能的观点,重视名词在和动词发生关系时所起的角色作用,但是怎样定义这些角色往往会有主观性,不便于实证。Bloomfield(1933)采取了实证的态度,尽可能不使用这些功能概念,但有时候不得不提到这些概念。从 Harris(1946)到 Chomsky(1965),采取的是实证性更强的形式语法的研究路子,力图取消这些功能层面的概念,即消除语法结构关系、语义结构关系,用更严格的初始概念来推导这些关系。Fillmore(1968)提出的格语法,实际上是重新回到语义结构关系,即施事受事这样一些关系,所以 Fillmore 本质上考虑的是组合中的语义功能,不过是在更高的层次上考虑问题。

格语法分三大部分:基础部分、转换部分和语音规则部分。后面两部分和 Chomsky(1965)标准理论(Standard Theory)基本是一致的,下面我们重点讨论和本章有关系的基础部分。

Fillmore(1968)在格语法理论中打算用格的分析来修正标准理论中的深层结构,他认为存在一个比深层结构更深的层面:语义格层面。Fillmore 的基础部分包括紧密联系的两个部分,一是改写规则,一是词典。Fillmore 提出了基础结构这个层面,提出句子的结构包括两大部分,下面分别加以讨论。

Fillmore 改写规则的基本格局是:

$$S \rightarrow M + P \text{ (Sentence} \rightarrow \text{Modality} + \text{Proposition)}$$

$$P \rightarrow V + C_1 + \dots + C_n$$

$$C \rightarrow K + NP$$

其中符号的含义是：

句子(S)：包含情态和命题

情态(M)：情态主要指动词的狭义情态、时、体、肯定否定等内容。

命题(P)：命题是动词和名词的各种关系，体现为施事格、受事格、工具格、与事格等。

C是由格标记K和名词短语NP构成，K可以是介词、前缀、后缀等。

Fillmore并没有给情态严格的定义。对他的用例进行分析可以看出，这里的情态是广义的情态，是指词典动词入句后的各种变化形式和附加形式。其实从句子中分出情态，Chomsky(1965)已经这样做了，Chomsky把句子分成：

$$S \rightarrow NP \wedge \text{Auxiliary} \wedge VP$$

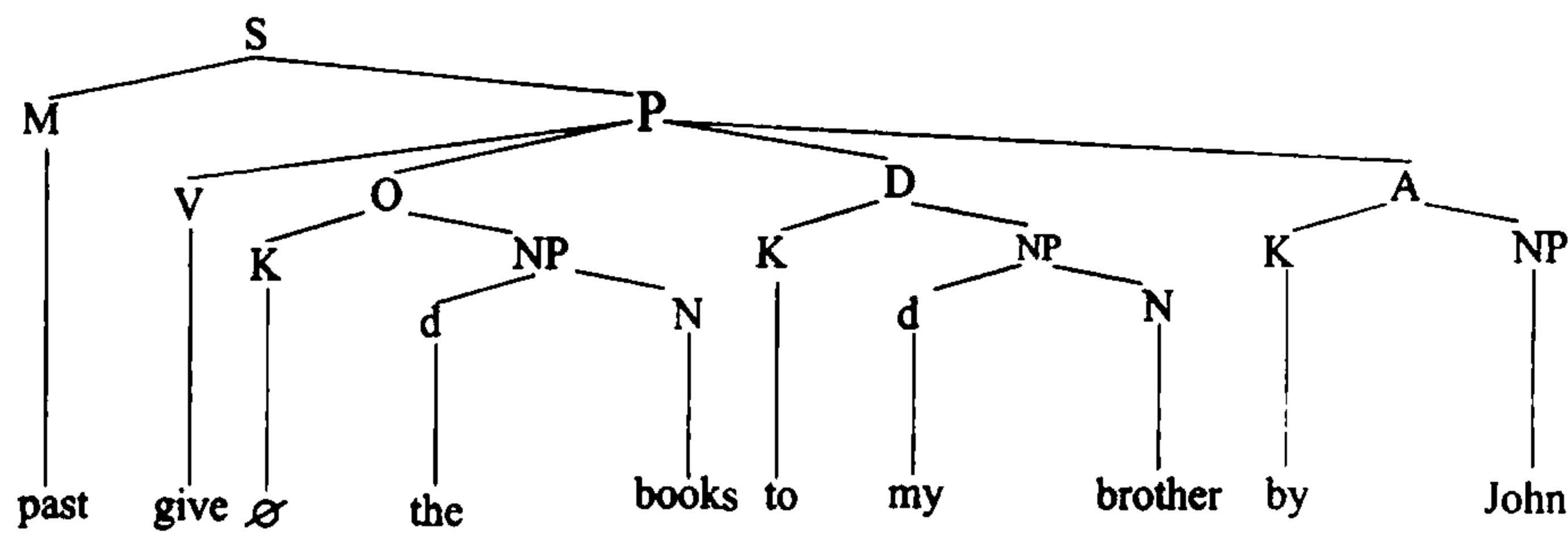
Fillmore的Modality相当于Chomsky的Auxiliary。Fillmore的改写规则和Chomsky(1965)标准理论的最大区别是：

1. $P \rightarrow V + C_1 + \dots + C_n$ 代替了 $VP \rightarrow V \wedge NP$ 。
2. P中的V和各种语义格Ci没有线性关系。

因此 Fillmore(1968)的基础部分的基本框架为：

基础→基本规则→S→M+P (M:情态;P:逻辑命题)
↓
└→ P→V+C₁+C₂+……+C_n(C:格成分)
底层格表

比如:



这个基础格式经过转换可以生成下面一系列形式:

- John gave the books to my brother
- John gave my brother the books
- The books were given to my brother by John
- My brother was given the books by John

语类范畴 V、NP 仍然是保留的,但短语结构规则已经被命题结构代替。语义格被进一步改写后才有 NP 的出现,NP 是比语义格更浅层的范畴。

Fillmore(1968)给出了几个基本的语义格,包括施事格(A=Agentive)、工具格(I=Instrumental)、与格(D=Dative)、使成格(F=Factitive)、处所格(L=Locative)、客体格(O=Objective)。后来,Fillmore 以及其他一些学者也对格作了扩展,下面把讨论得比较多的格归纳如下(以英语为例):

格分类	简写	格名称	格标记举例	功能
Agentive (or Agent or Actor)	A	施事格	by	动作的发出者, 通常是有生命的
Objective	O	客体格	零	动作的接受者
Theme (or Patient)	P	受事格	零	行为的接受者
Instrument	I	工具格	with, by	动作或行为的工具
Factitive		使役格	to	动作行为的终点
Locative	L	处所格	in, at, on, to, from……	动词确定的动作或状态的处所或方向, 格标记的选择有的由动词决定, 有的由名词决定
Time	T	时间格	in, at, on, to, from, during……	事件发生的时间。格标记的选择有的由动词决定, 有的由名词决定
Dative	D	与格	to	动词确定的动作或状态所影响的有生生物
Benefactive	B	受益格	for	行为的受益者
Comitative	C	伴随格		动作行为的伴随者
Experience	E	经验格		行为经验者
Source	S	来源格	from	行为的起因
Goal	G	目的格	to	行为的目的

比如：

- [John^A] opened the [door^O]
- [The door^O] was opened by [John^O]
- [The key^I] opened [the door^O]
- [John^A] opened [the door^O]with [the key^I]
- [John^A] used [the key^I] to open [the door^O]
- [John](D) believed [that he would win^P]
- [We^A] persuaded [John^D][that he would win^P]
- It was apparent to [John^D] [that he would win^P]
- [Chicago^I] is windy
- It is windy in [Chicago^I]
- There are [many toys] in [the box^I]
- [Mary^T] fell over
- [Mary^F] was happy
- [John^A] sang a song for [Mary^B]
- [John^A] passed the book to [Mary^G]
- [John^A] returned [from paris^S]

到底需要多少语义格，标准是什么，目前还在讨论中。Fillmore 认为语义格是人类语言的普遍现象，我们认为，可以承认语义格这个层面在人类语言中是普遍的，但在具体语言中，语义格的分类就像词类一样，在不同的语言中并不一样。一个语言到底需要多少种语义格，还需要深入研究。

汉语中的介词有时候也可以用于判定语义格：

施事	_V, 被 NP	杯子被[汪锋]砸了	
受事	V_, 把 NP	汪锋把[杯子]砸了	
与事	V_NP, 给 NP	汪锋给[邵永海][一瓶云南咖啡]	
对象	对于, 对	美国对[中国]施压	

语义格并不限于名词,上面已经给出英语中命题作语义格的。汉语中也有命题作语义格的。比如:

[小邵^{施事}]告诉[汪锋^{与事}][郭锐的儿子很漂亮^{命题}]

Fillmore 主张区分名词词典和动词词典。名词词典至少应该有下列的信息:

- 名词词典:a. 名词可能作什么样的格,如 A,O,I
b. 名词固有的语义特征,比如“有生命、无生命”

在动词词典中,动词除了通过格框架相互区别开,还通过转换特征(transformational properties)区别开。转换特征包括 3 个方面:

The most important variables here include (a) the choice of the particular NP to become the surface subject, or the surface object, wherever these choices are not determined by a general rule; (b) the choice of prepositions to go with each case element, where these are determined by idiosyncratic properties of the verb rather than by a general rule; and (c) other special transformational features, such as, for verb taking S complements, the choice of specific complementizers (that, -ing, for, to, and so forth) and the later transformational treatment of these elements.

由于这里谈论的 3 个方面包括语义格作主语的情况、介词选择的情况、附加成分等,这些方面尽管都是和转换有关系的,

但不是普遍规则,是词典信息,所以要放到词典中处理。

Fillmore 认为动词和语义格搭配的时候往往有固定的介词引入,所以动词词典中要给出格的介词标记清单:to, in, on, at。例如:

good at
arrive in

这就是前面给出的 $C \rightarrow K + NP$ 的要求。其实哪些格需要引入介词,引入什么样的介词,是有区别的。比如:

He was good at cooking
He died in Newyork
He died in 1946

第 1 个例子是固定搭配,动词词典中需要提供介词搭配信息。第 2 个例子并不是固定搭配,只要是在特定地点可能发生的行为动词,都可以用 in 引入处所。第 3 个例子中的介词也不是动词 die 的搭配属性,并不需要词典提供信息,而是英语中表示年份必须使用的介词,也不应该在动词词典中提供信息。这方面的讨论我们后面还要展开。

Fillmore 的格语法与 Tesnière 的配价理论有一定的关系,Fillmore(1968)在他的研究中也提到了 Tesnière 的工作。Tesnière 于 1939 年首次阐释了从属关系语法的基本论点。1959 年的《结构句法基础》又详细展开了他的思想。从属关系语法最基本的概念是关联(connexion)和转位(translation)。法语 Alfred parle(Alfred 讲话)中,Alfred 和 parle 之间有一种关联关系。关联是有层次的,这就是从属关系。在 Alfred mange une pomme(Alfred 吃苹果)中,动词 mange(吃)是句子的“结”

(noeud), Alfred 和 pomme 从属于动词 mange, une 从属与 pomme。Tesnière 认为动词是句子的中心, 宾语和主语受动词支配, 宾语和主语是平等的。与此相关的一个重要概念是动词的“价”。与动词直接发生关系的有名词词组构成的“行动元”(actant)和由副词词组形成的“状态元”(circonstants)。行动元的数目不得超过三个, 即主语、宾语 1、宾语 2。行动元的数目决定了动词的价(valence)。比如:

零价动词(verb^{es} avals):

il pleut (下雨)

一价动词(verb^{es} monovalents):

il dort (他睡觉)

二价动词(verb^{es} bivalents):

il mange une pomme (他吃苹果)

三价动词(verb^{es} trivalents):

il donne son livre à Charles (他把他的书给 Charles)

状态元从理论上说是无限的。Tesnière 主要是从句子成分来解释动词的价, Chomsky 在解释动词框架的时候从句子成分入手, 思路和 Tesnière 有相似的地方。从语义结构关系上看, Fillmore 的语义格和动词的价也有共同的地方。从统一的角度看, 动词的“价”是由动词的“格”框架决定的, 所以“格”是最基本的、更初始的概念, 决定了格就决定了价, 但决定了价不一定就决定了格。如“一价”的动词可以有“人来了”, 也可以有“天亮了”, 但“人”和“天”在这里的语义格是不一样的, “人”是施事, 而“天”不是。

语义格有很强的解释力。如果动词入句后语义格不完备,就会产生歧义:

- a. 来的是[小孩^{施事}]
- b. 看望的是[小孩^{施事、受事}]
 看望母亲的是[小孩^{施事}]
 母亲看望的是[小孩^{受事}]
- c. 送的是[小孩^{施事、受事、与事}]
 送花的是[小孩^{施事、与事}]
 送医院的是[小孩^{施事、受事}]
 我送的是[小孩^{与事、受事}]
 送医院鲜花的是[小孩^{施事}]
 我送医院的是[小孩^{受事}]
 我送(他)鲜花的是[小孩^{与事}]

右方括号左上方的名词标注了“小孩”可能作哪些语义格,如果可以作两个或两个以上的语义格,就形成了歧义,只有当句子中语义格得到补充,歧义才可能消除。这里的歧义也可以用“价”的数量来解释。“来”等是一价动词,“看望”等是二价动词。“送”是三价动词。显然,价只是给出了可能形成歧义的数量,而语义格的标注可以让人看出歧义的形成是缺少哪些格引起的。

动词和语义格的关系后来得到广泛研究,形成一套题元理论(θ -theory)。动词和语义格的关系又被称为题元关系(Thematic relation)或论元关系(Argument relation)。题元关系也被称为题元结构(theta structure),论元关系也被称为论元结构(argument structure)。动词的语义格具体项目也被称为题元角色(thematic role, 又被称为 theta-role, θ -role)。argument 最早是逻辑学中的术语(Carnap 1937),指谓词(predicate)的参与者,也翻译成变元或主目。从文献材料看,论元(变元、主目)、

语义格、题元角色在后来的工作中逐渐形成了细微分工。比如，从 Chomsky(1981, p36) 给出了的题元准则(θ -criterion)可以看出论元和题元角色的区别：

Each argument (论元) bears one and only one θ -role, and each θ -role is assigned to one and only one argument.

从具体用法看，论元多指深层结构或线性结构中有一定位置的语义格、及其数量，和 Tesnière 的“价”的概念更接近，题元角色多指取得特定语义角色(如施事、受事)的语义格。通过对下面实例的解释可以有更具体的理解：

[房子^{客事}]垮了

垮：一价动词，论元结构一价，一元谓词

[家长^{施事}]拜访[老师^{受事}]

拜访：二价动词，论元结构二价，二元谓词

[老师^{施事}]给[家长^{与事}][一封信^{受事}]

给：三价动词、论元结构三价，三元谓词

[邵永海^{施事}]告诉[汪锋^{与事}][郭锐有一个漂亮的儿子^{命题}]

告诉：三价动词，论元结构三价，三元谓词

以上方括号中的成分都是语义格或论元或题元，右括号左上方则是具体的角色。

由于语义格最初是相对于表层格引入的，表层格多讨论表层名词格的性质，语义格也多涉及语义结构中名词的性质。实

际上语义结构关系也涉及动词周围谓词性卫星成分,比如“我听说[小邵感冒了]”这样的句子。Fillmore(1968)也提到语义格可以是句子。后面我们更多的使用论元结构和论元方面的术语,以避免把语义格理解成和名词有关系的成分。

8.2 语义格和线性深层结构的关系

Gruber(1965)、Fillmore(1968)、Jackendoff(1972)主张,句子的语义结构应该是语义格关系,其中尤其以 Fillmore(1968)的研究最有代表性。Fillmore(1968)陈述了他的功能谓词

个动词的施事、受事、与事、工具等才有可能提升为主语。不过 Fillmore 没有论证这种主语提升过程是否可以通过次范畴化得解释。

Fillmore 也没有讨论为什么语义格普遍存在的理由。这问题是语义结构关系的初始性问题。下面我们将论证“语义格”个层面应该是普遍存在的,但各种语言中的语义格可能并不完全一致的,这是因为不同的语言动词和动词之间没有一一对应关系。如英语 please:

The children's performance pleased the audience

这样的动词是不能用汉语对译的。因此,语义格的初始性问题并不是说人类所有语言的语义格分类都是一样的,并且每种语言必须具有这些语义格。语义格的初始性问题是说,在人类语言中存在一个语义格的层面,这个层面不能用语类、直接成分、语法关系来取代。

语义格面对的问题本质上就是语义解释的问题。Katz 和 Postal(1964)、Chomsky(1965)认为句子的深层结构决定了语义解释。前面提到,Chomsky 在处理语义时,采取的方法是对语类的次范畴化,对动词作选择性规则的限制,所以 Chomsky 的做法是句法中心论,他的深层结构是句法深层结构。NP、VP 等语类范畴都是深层结构中的根本范畴,并且有固定的顺序。在句法深层结构中,并没有更多的考虑语义结构关系,即没有考虑“施事、受事、与事、目的、工具”等功能概念。从根本上说,没有承认组合关系的两大部分:语法组合关系和语义组合关系。不过从另一个角度说,标准理论已经涉及语义组合关系,是在次范畴化中涉及的。按照 Chomsky(1972)的意见和对 Fillmore 格语法的回应,语义格的信息已经蕴含在动词次范畴中。概括地说,Chomsky(1965)的讨论不涉及施事、受事等语义功能观念,只到语类和语类的位置为止。为了进一步弄清楚次范畴化和语义格到底是什么关系,我们需要严格区别句法功能(或语法功能)和语义功能。主语、谓语、宾语等是句法功能,施事、受事等是语义功能。Chomsky(1965)认为语类以及语类顺序已经蕴含了句法功能。

从实证的角度看,Chomsky(1965)是在深层结构部分坚持线性化过程,即通过语类的线性分布来解释各种句法功能问题。这样一种线性化过程的好处是具有可观察性。其实从 Bloomfield(1926)早期分布理论到 Harris(1946、1951)的发现程序,涉及的都是—种线性分布,所以也可以看成是—种线性化过程,只

不过 Bloomfield 和 Harris 的线性观念是指可观察到的语类序列的线性观念。Chosmky 更明确认为短语结构规则蕴含有功能,并且把这一性质限制在深层结构。格语法和 Chomsky 的次范畴处理方式并不同,语义格是定性的,功能取向的,即需要给出或定义“受事、施事”等功能角色,但不需要限定线性顺序。格语法定性的思想比较直观,容易接受。但是,现实语言中语义结构关系太复杂,有些动词和名词之间,是什么语义结构关系,现在还说不清楚。Chomsky(1965)不直接提语义结构关系,不直接讨论施事、受事等概念,也有一定的好处,即便于形式化。用次范畴化和选择规则来控制语义关系,可以避免对施事、受事等语义格的定义和定性问题,类似结构主义用分布和词类来控制语法结构关系,可以避免语法结构关系本身的定义和定性问题。

考虑下面 Chomsky(1972)讨论过的一个实例:

John /opened the door
the key /opened the door

按照 Chomsky(1965)的标准理论,这两个结构都是“主语—谓语”结构:

主语	谓语
John	opened the door
the key	opened the door

按照 Fillmore 的解释,这两个结构是不同的,前一个是“施事-行为”结构,后一个“工具 行为”结构:

施事	行为
John	opened the door. •

工具	行为
the key	opened the door

Fillmore(1968)的倾向性意见是,只需要语义格模型解释语义,可以不要深层结构。我们前面已经提到,Chomsky(1972)认为格语法和深层结构理论没有根本的区别,格语法关于语义结构的区别已经蕴含在标准理论(Standard Theory, Chomsky, 1965)的深层结构中。我们认为这两个层面既有共同点,也有区别。拿上面的实例看,有生命的主语是施事,没有生命的主语是工具,而标准理论是区别有生和无生这一对范畴的,因此,格语法和标准理论在识别施事、受事、工具上没有根本区别,一些学者关于语义格存在的理由我们认为都可以通过标准理论中建立在次范畴化基础上的深层结构加以解释。

其实在很多情况下,标准理论和格语法确实是等价的。考虑下面两个深层结构或线性结构:

猫咬狗
狗咬猫

这两个结构都可以转换出下面的句子:

咬的是猫

这是大家都熟悉的一个歧义句式。从转换的机制看,我们有两种转换,其表层形式是相同的:

NP1 + V + NP2
→ V — 的 — 是 — NP1
→ V — 的 — 是 — NP2
•

这是歧义的原因所在。按照 Chomsky(1965)线性深层结构的语义解释理论来看,“咬的是猫”中的 NP1,也可以是 NP2,即既可以是深层结构中的主语,即动词前的 NP,也可以是深层结构中的宾语,即动词后的 NP,所以有歧义。按照 Fillmore(1968)的理论来解释,“咬的是猫”中的“猫”可以是施事,也可以是受事,所以产生歧义。

再来分析一下“咬”的词库结构特征。用线性深层结构描写是:

咬:NP(动物)+V+NP

如果用格语法描写是:

咬:V {施事(NP,动物);受事(NP)}

这里的 V 所带的语义格没有线性位置。以上两种描写都能判定下面的句子没有歧义:

咬的是书

因为无论是线性深层结构中的主语,还是格语法中的施事,都要求有“动物”的语义特征,“咬的是书”中的“书”由于没有“动物”特征,所以只能理解成线性深层结构中的主语或深层格关系中的施事。所以两者在解释以上歧义方面是等价的。两者在很多方面的语义解释也和这里的情况类似,是等价的。

就上面的实例看,格语法和深层结构的区别在于,格语法对“施事—行为、工具—行为”的解释是直接的,标准理论的深层结构对“施事—行为、工具—行为”的解释是间接的。如果重视这种区别,深层结构语义解释理论和格语法语义解释理论代表了两种不同的语义解释模型,深层结构语义解释理论是句法中心

论,格语法的语义解释模型是语义功能中心论。

8.3 语义格的初始性

我们先论证早期一些学者关于语义格初始性论证的不充分性。我们重点分析 Radford(1988, p372)给出的一些典型英语例子。首先看一下作格例子:

John rolled *the ball* down the hill

The ball rolled down the hill

无论 *the ball* 在第一个句子中作宾语(VP 的 NP)还是在第二个句子中作主语(S 的 NP),语义结构作用都是相同的,都是 theme(受事),后一个句子就是学者们通常所说的作格结构(ergative structure),这种结构中的动词称为作格动词,即及物动词的宾语与不及物动词的主语有相同形式(即同格)或者说及物动词的宾语在该动词的不及物用法中作为主语使用(Perlmutter, 1978; Burzio, 1986)。这两个句子中的 NP(*the ball*)的选择关系是一样的:

(a) John rolled *the ball* / *the rock* / ! *the theory* / ! *sincerity*
down the hill

(b) *The ball* / *the/rock* / ! *the theory* / ! *sincerity* rolled
down the hill

由于汉语没有强式被动标记,所谓作格结构的用法在汉语中比较普遍:

转轮子

轮子转了

滚皮球

皮球滚了

开门

门开了

先不论作格的理论是否合理,这里的问题的实质是,动词的某个论元可以出现在句子的主语和宾语位置上,语义格或论元的角色并没有改变,Chomsky(1965)的次范畴化和选择信息似乎没有解释这类现象,因为在 Chomsky(1965)的模型中,NP 的位置是固定的,不同位置的 NP 有不同的语义解释。不过,Chomsky 的线性模型可以在词典中列两个 roll,在主动句中这两个 roll 前面的主语的语义特征可以不同。这种方案是允许的,因为英语中哪些动词属于作格动词,并没有找到周遍条件,所以作格动词实际上也必须在词库中标注。从这个意义上说,作格动词的存在并不构成语义格层面初始性的理由。

再看语义格层面初始性的另一种可能理由。有些名词尽管都是同样作主语,或者说作 S 的 NP,但语义结构作用是不同的:

The bottle shattered the glass (瓶子打碎了玻璃)

The bottle shattered (瓶子打碎了)

NP(the bottle)的选择关系也是不一样的:

(a) *the vase /! the noise /a hidden flaw shattered the glass*

(b) *the vase* /! *the noise* /! *a hidden flaw* shattered

汉语也有类似的例子:

傅刚伤了汪锋的脚

傅刚伤了脚

王冕死了父亲

王冕死了

这也是句子成分、范畴信息、次范畴信息以及选择信息没有讨论的问题。从这个角度看,语义格似乎能够向我们提供句法结构分析或语类结构分析所不能提供的相似性和差异。不过次范畴化理论也可以说上面的 VP 本来就不一样,所以对前面的 NP 的选择也不一样。

Radford (1988, p375) 同意 Fillmore (1972) 和 Gruber (1976) 的主张,副词的共现规则和论元角色有关系。Gruber (1976) 认为,像 *deliberately* 这样的副词只能和语义格 Agent 短语共现:

(a) John (= AGENT) deliberately rolled the ball down the hill

(b) * the ball (= theme) deliberately rolled down the hill

不过这样的实例还不能证明语义格的初始性,因为根据“有生”这一特征也可以解释这里的现象,即 *deliberately* 只能和“有生”的深层主语共现。

Fillmore (1972, p10) 认为, *personally* 这样的副词只能和 Experience 语义格(主目)共现:

- (a) Personally, *I* (= EXPERIENCER) don't like roses
 (b) Personally, your proposal doesn't interest *me* (= EXPERIENCER)
 (c) * Personally, *I* (= AGENT) hit you
 (d) * Personally, you hit *me* (= THEME)

其实这个实例也不能证明语义格的初始性。这里 personally 的出现条件应该和动词有关系, like 是允许说话人有评价态度的, hit 不允许。这实际上是一个普遍的认知规律, 在汉语中也存在类似情况。下面是允许有评价态度的动词:

就我个人来说,我是想去的
 就我个人来说,我是喜欢她的
 就我个人来说,是愿意在上海的
 就我个人来说,我是宁愿学数学的
 就我个人来说,我是希望当老师的

下面是不允许有评价态度的动词:

* 就我个人来说,我打了他
 * 就我个人来说,我救了他
 * 就我个人来说,我买过书
 * 就我个人来说,我知道他
 * 就我个人来说,我是教师

从更一般的角度看,以上的区别和谓语是否允许评价有关系,而不仅仅是动词的问题,比如“喜欢”这样的允许有评价态度的动词,如果用在已经发生的事情上,就不允许评价了:

* 就我个人来说,我喜欢过他

* 就我个人来说,我同意过他

而不允许有评价限定的动词,如果加上情态,也可以和评价类词语共现:

就我个人来说,我可以打他的

就我个人来说,我可以救他的

就我个人来说,我应该买书的

就我个人来说,我可以当老师的

所以这里的共现问题,仍然是认知搭配问题,并不是语义格作为初始概念的语言学证据。

Anderson (1977, p267-271)认为,在英语中,只有带施事主语的动词才可以有带 *by phrase* 的配对名词:

(a) The *mayor* (= AGENT) protested

(b) the protest *by the mayor*

(a) The *mayor* (= THEME) died

(b) * the death *by the mayor*

不过这里的实例也可以用及物动词和不及物动词来解释。对于:

1. NP1—V—NP2

2. NP1—V

只有满足 1 的动词,才有带 *by phrase* 的名词配对。所以 Anderson 的实例还不能证明语义格的初始性。

Kayser 和 Roeper (1984)、Roberts (1985)等注意到英语中

有一种现象：

The book reads easily

The paper writes difficultly

这种现象被称为中动现象，即及物动词的受事作主语，动词并不是被动态。这种句子被称为中动句(Middle Structure)。中动问题的研究有另外的动机，后来的研究也有很多有价值的解释。我们在这里提出中动句来讨论，是因为这种现象似乎违背了及物动词的选择限制，即 read、write 这样的动词前面本来应该是有生的 NP。但中动句仍然不是格语法初始性的理由，因为中动句是有条件的，在语义上表示 NP(book, paper)的一种性质，在形式上有表示性质的副词出现在动词后面。这种形式和语义对应就可以通过选择限制来控制，即只有当这样的形式出现的时候，非有生的 NP 主语才有这样的用法。

Fillmore(1968,3)认为句子连接的条件不是句子成分相同而是语义格相同：

Only noun phrases representing the same case may be conjoined.

考虑下面的情况：

18. Jonh broke the window

19. A hammer broke the window

20. John broke the window with a hammer

以上句子的连接情况如下：

Jonh and Jack broke the window

* Jonh and a hammer broke the window

Fillmore 认为后一句违背了语义格相同的原则。Fillmore 想说明连接原则也能证明语义结构关系是初始的。不过,我们仍然可以用普遍认知原则来解释这里的现象,由于 John 和 hammer 有不同的语义特征,John 是“有生”的,“hammer”是“无生”的,所以不能形成并列关系。这类情况在汉语中也是存在的,而且不限于主语位置:

小邵打破了窗子

榔头打破了窗子

* 小邵和榔头打破的窗子

小邵吃面条

小邵吃食堂

* 小邵吃面条和食堂

小邵写毛笔

小邵写信

* 小邵写毛笔和信

小邵挖煤

小邵挖洞

* 小邵挖煤和洞

从 Chomsky(1965)的次范畴角度看,也可以说是选择关系不一致。一般地说,线性深层结构中充任主语的几个名词应该有相同的选择限制。

以上给出的理由说明语义格在解释语义和句法的时候有很重要的价值,但这些理由本身并不证明语义格的初始性或非推导性,这些语义格理论所解释的问题根本上都可以从此前我们讨论的次范畴化的概念中得到解释,只是在有些方面用次范畴化解释起来比较复杂。Chomsky(1981)和很多学者后来都接受了语义格的概念,即引入了题元的概念,但并没有论证题元作为初始概念存在的理由。因此,引入语义格这个层面,到底是因为研究的方便还是语义格层面具有初始性,仍然没有得到证明。

现在我们从另一个角度来考虑语义格的初始性。问题仍然是,语类在深层结构中的线性分布是否一定像 Chomsky(1965)所给出的分析那样,蕴含了语义功能,即蕴含了施事、受事等功能? 让我们来分析下面的结构:

NP1 V NP2

根据 NP 的位置、语义特征或次范畴化理论,是否一定能够推导出 NP1 是施事, NP2 是受事?

在汉语中首先遇到的问题是 NP 位置互换但不改变语义结构关系:

一张桌子坐两个人

两个人坐一张桌子

因此要用 NP 的位置来确定施事有一定的困难。汉语中遇到的另一个问题是,有些情况下次范畴解释语义格确实会遇到困难。考虑下面实例:

写字

写了字

写好字

写钢笔 * 写了钢笔 * 写好钢笔

这两种转换是不一样的,必须要有“结果、工具”这样的语义格才能说清楚。如果要坚持说这里的“字、钢笔”在次范畴上有什么区别就有一定的困难。这就提出了一个根本性的问题,次范畴化到底需要细化到多深?最深的底线是任何两个词都是不同的类或次范畴。最能证明次范畴化不能替代语义格的根据是,同一个词也可能有不同的语义格,这就突破了次范畴化的最低限度。考虑下面的实例:

吃碗 把碗吃了 用碗吃
吃筷子 把筷子吃了 用筷子吃

以上“碗、筷子”都有两种理解,一种是工具,一种是受事。同样一个词有两种含义,已经不能用次范畴来解决问题了。歧义的产生不是因为碗、筷子的语义特征,而是因为可以有工具格和受事格两种情况。也就是说,即使我们分出“可食”和“不可食”两种语义特征,说在“可食”的语义特征下是受事,在“不可食”的语义特征下是工具,也不能消除这里歧义,因为碗、筷子在特定的语境下仍然是可食的,即使是可食的,也可以作工具。有没有歧义主要在于把筷子理解成工具还是受事。汉语中这类例子是很多的。

其实用语义特征来进行次范畴化,进而控制语义结构关系,这种方法并不能贯彻到底。比如用“可食”来限制“吃”的受事范围,本身就存在问题。考虑下面实例:

吃石头/吃沙子/吃泥土/吃空气/吃地球

很难说“石头、沙子、泥土、空缺、地球”有可食的特征,但在一定

的条件下上面的组合都可以说。比如说,一个嘴巴地球还要大的外星人“吃地球”是合法的句子。

再考虑更复杂的动词:

小邵说服汪锋去上海

小邵答应汪锋去上海

两个句子的次范畴化规则是相同的:

NP+V+NP+V

(_N小邵)(_{VP}[_V说服][_N汪锋][去上海])

(_N小邵)(_{VP}[_V答应][_N汪锋][去上海])

但语义结构关系并不一样:

(_N小邵^{施事}说服)(_{VP}[_V说服][_{VP}汪锋^{受事·说服、施事}去][去上海])

(_N小邵^{施事/答应}去)(_{VP}[_V答应][_{VP}汪锋^{受事/答应}][去上海])

在“说服”的语义结构中,前面的 NP 是施事,后面的 NP 是受事,同时是另一个补语动词的施事。在“答应”的语义结构中,“答应”前面的 NP 是 V 的施事,也是补语动词“去”的施事。“答应”后面的 NP 只是“答应”的受事,和后面补语动词“去”没有语义格关系。以上的次范畴框架或短语标记(phrase marker)是一样的,语义结构关系并不相同。这种语义结构关系上的区别也不可能通过名词的次范畴化解决,因为上面两个 NP 在两个句子中都是相同的。

以上论证了语义格或论元结构的初始性,由于论元并没有固定的线性位置,因此论元结构是非线性的。“施事、受事、与事、工具”等语义格是从语义角色或语义功能定义断定单位在句

子中的性质,因此,语义格的初始性也进一步说明语义功能概念在基础部分具有初始性。也就是,人类语言的基础部分不能不考虑语言功能。在后面的相关章节,我们还会从句法结构关系论证功能的初始性。现在我们可以说次范畴并不能完全解释语义格问题。语义格层面是一个初始层面。当然,语义格这个层面的存在并不能否认与次范畴分析相关的语义特征分析,“吃碗”之所以允许有“把碗吃了”的用法,是因为“碗”在特定的语境下具有作为“吃”的受事的语义特征。比较:

吃面

吃大碗(有歧义)

? 吃绝对

“吃绝对”不合法,这说明即使把“绝对”放在受事的位置,“绝对”的语义特征也不能实现和“吃”的组合关系。

8.4 词库中动词语义格的确定和直接关联原则

语义格或语义角色或论元到底有多少? Fillmore(1968)给出的几种语义格是相当概括的。后人作了更细的区分。但是,如果没有一个分类的目的和标准,更细的语义格分类是否能够解释语义结构关系的本质? 这种细致的区分如果没有一个有效的标准,就会使问题变得琐碎、盲目而没有应用和认识价值。

要给出语义格的判定标准,涉及对语义格的本质的进一步认识。其实 Fillmore(1968)的理论已经暗含了一些判定标准,只是还不明确,并且有些地方有矛盾。Fillmore(1968)认为一种格关系只能在简单句中出现一次。这是语义格的唯一性原则。下面的例句中,一种格关系出现了两次,是不合法的:

' A hammer broke the glass with a chisel

由于 chisel 是无生命的 (inanimate), 因此这个名词不能理解成施事, 只能理解成工具。又由于 hammer 和 chisel 都只能理解成工具, 这个句子违背了语义格的唯一性原则。这一原则在人类语言中也有普遍性。

根据唯一性原则, 我们立刻可以推演出一个结论: 如果在简单句中两个不同位置的 NP 直接和 V 发生关系并且这个简单句是可接受的, 这两个 NP 一定是两个不同的语义格或论元。更进一步, 如果在简单句中 n 个位置不同的 NP 直接和 V 发生关系并且这个简单句是可接受的, 这 n 个 NP 一定是 n 个不同的论元。这也可以看成是一个基本的原则, 我们把这一原则称为语义格或论元的共生原则。共生原则是确定语义格的一个基本原则。

郭锐 (1995, p170) 在研究汉语句法时也注意到这类现象, 认为如果不同语义角色可以在同一小句的主宾语位置上共现, 则属于不同论元角色; 如果总是不在同一小句的主宾语位置上共现, 则可合并为同一论元角色。郭锐意识到的共现问题非常重要。当然, 我们可以不把判定的基础限制在主宾语位置, 因为表层主宾语和动词的关系比较复杂, 把论元的共现范围限制在主宾语, 会遇到下面的问题:

这张桌子我砍了腿

这个苹果我削了皮

这个问题我找出了原因

其中有些句子成分不一定直接和动词发生语义结构关系。上面三个句子的主语可以看成主谓谓语句的主语而不是动词的直接语义论元:

[这张桌子][我砍了腿]

[这个苹果][我削了皮]

[这个问题][我找出了原因]

从另一个角度看,这些句子又可以看成是下面句子的转换结果:

我砍了[这张桌子]的腿→[这张桌子]_i我砍了 t_i腿

我削了[这个苹果]的皮→[这个苹果]_i我削了 t_i皮

我找出了[这个问题]的原因→[这个问题]_i我找出了 t_i原因

即主语是受事宾语的定语 NP 移位并删除该定语后面“的”的结果。上面的 t_i 表示移位后的语迹,下标 i 表示同指,即指称相同。根据表层结构确定论元还会有其他的问题,这里不展开。为了避开表层句子成分的干扰,确定语义格或论元需要从语义关系的直接关联入手。直接关联将不限于讨论主宾语位置上的语义格。只要直接和 V 共同发生语义结构关系的,包括由介词引入的成分,都可以看成是语义格。基于这一点,我们可以给出一个判定动词的语义格或论元的直接关联原则:

一个动词的语义格是可以直接和动词同时关联的全部实词短语。

在有些情况下,不同性质的 NP 可以出现在同一个句法位置上:

我吃筷子

我吃饭

我吃食堂

这些 NP 算几个语义格? Fillmore(1968,3)还提出了一个并列连接程序来确定语义格的同一性,即认为只有代表相同语义格的 NP 才能并列连接(Only noun phrases representing the same case may be conjoined)。根据 Fillmore 的这一观点,上面“吃”后面的 NP 都不是相同的语义格,因为不能并列连接。这一分析结果和很多汉语研究学者直接靠语义分析得出的结果是一致的:

- * 我吃筷子和饭
- * 我吃筷子和食堂
- * 我吃饭和食堂
- * 我吃筷子、饭和食堂

但如果考察更多的实例,Fillmore 的并列连接原则会遇到困难,下面的 NP 也相互替换:

- 他在挖地
- 他在挖洞

并且也不能并连:

- * 他在挖地和洞

按照 Fillmore 的原则,“地”(受事)和“洞”(结果)也应该看成是不同的语义格。但是“洞”和“地”显然不能在以“挖”为中心的简单句中 和动词“挖”直接关联。可见,Fillmore 并列连接原则和我们给出的直接关联原则是有矛盾的。如果坚持并列连接原则放弃直接关联原则来提取论元,论元的种类会相当复杂。如果坚持直接关联原则而放弃并列连接原则,怎么解释“地、洞”不

可连接的问题？我们认为不可并联并不是受论元不同的限制，而是受更严格的语义分类或次范畴化的限制，因此论元的确定主要应该依据直接关联原则而不是次范畴化原则。这也许从另一个侧面证明论元结构和次范畴化（或动词结构框架）不完全一样。

仍然以上面的实例为例。“吃”的“筷子”（工具）、“饭”（受事）和“食堂”（处所）三种语义格可以在不同的位置上共存：

我在食堂用筷子吃饭

所以对“吃”来说，“筷子、饭、食堂”是不同的语义格。而“挖”所带的成分“地”（受事）和“洞”（结果）不能在不同的位置上和动词直接关联，应该是相同论元的变体。“在地上挖洞”是另一种情况，“地”是处所。

Fillmore(1968)认为不同的语义格不能并联，现在看来，不仅不同的语义格不能并联，同一语义格的不同变体也不能并联：

挖水沟、窑洞（都有受事特征）

* 挖树根、水沟（类型并不相同）

因此，NP 并联的限制条件不仅需要相同的语类，也需要相同的语义特征和语义格。

根据直接关联原则，很多起不同功能作用的名词都可以归并成有限的语义格。比如，工具名词和方法名词不能和动词发生直接关联：

我用节能办法烧水

我用电饭煲烧水

这个电饭煲我烧水

这个节能办法我用来烧水

* 这个电饭煲我用节能办法烧水

· 这种节能办法我用电饭煲烧水

因此工具和方法应该作为同一语义格的变体。

Fillmore 的并连原则不能有效区分不同的语义格和同一语义格的不同变体这两种情况。现在我们把直接关联原则作一个比较严格的定义：

同一语义格只能和一个动词进行一次直接关联，不同的语义格必须有资格同时和一个动词直接关联。

直接关联原则为语义格的不同角色划分提供了一个底线，即找不到共同直接关联语境的两 NP，不必分成不同的语义格。这些 NP 的语义差异另有作用，需要有其他办法来处理。至少我们需要做出两种区分。语义格是功能性的，是在组合关系中形成的功能，而次范畴化是一种分类，涉及的问题是语义搭配的限制条件。这两个层面的区别可以进一步从下面的实例得到理解：

吃鸡/吃面/吃肉/吃苹果

喝水/喝牛奶/喝果汁/喝酒

其实两组实例的直接关联关系都是动词和受事的关系，但由于选择的动词不同，名词也不同。

根据直接关联原则，我们可以区分出以下主要的语义格：

论元性质	论元	论元次范畴	介词或准介词
黏着语义格	主事	施事(agent)	被
		致事(causer)	让
		经事(experiencer)	
		感事(sentient)	
		主事(theme)	
	客事	受事(patient)	把
		结果(result)	把
		对象(target)	对、对于
		系事(relative)	
		受益(benefactive)	为
	与事	与事(dative)	给
		凭借	用、凭借、根据
		工具(instrument)	用、凭借
		方式(manner)	用、凭借
		方式(manner)	用、凭借
自由语义格	时间	时间点(time)	在
		起点(source)	从
		终点(goal)	到
	时段	时段(duration)	
	处所	处所(locative)	在
		起点(source)	从
		终点(goal)	到
		路径(path)	沿着,过
		路径(path)	沿着,过

袁毓林(2002)曾给出个一个论元层阶体系,和我们根据直接关联原则的区分不完全相同,但也是一种有价值的思路,可供参考。

特定的语义格及其变体往往有一组介词引导。Fillmore (1968)就尝试从形态格和介词的对应来确定语义格。这方面仍

然还有很多地方不清楚,但大致的对应是有的。我们在上面给出了汉语的介词和准介词对应语义格的情况,仅作参考。

有人可能会认为一个介词往往可以引介多种论元角色,因此反对用介词引介语义格。比如汉语的“给”:

给父母请安
给孩子做示范
给房子刷漆
给大家介绍
给新手提干
给朋友还钱
给运动员加油

如果从次范畴细分,“给”可以引导很多论元角色,实际上这些不同的论元角色按照直接关联原则都是不能和同一个动词同时直接关联的,这些不同的次范畴只是语义格变体,都可以统一在语义格“与事”下面。

我们认为,从语法标记(包括介词)的角度研究论元的性质有两方面的价值。首先,句法单位或短语的语义格身份不可能在词库中标注,只能在和动词组合的时候认定,这时介词等语法标记对判定语义格有关键作用。在汉语中,语义格往往通过上下文确定。但如果上下文有歧义,介词等语法标记就必须使用。从介词或语法标记的角度研究语义格的另一价值在于,一个语言关心哪些语义格,可以从形式上得到说明。

以上的语义格是非常粗略的区分。这里的自由语义格是指可以通过周遍规则描述的部分。有些时间、时段或处所词,和动词有固定搭配,比如“躺小床”,这时的“小床”应该属于客事。还有一些语义成分,如“原因(reason)、目的(aim)、范围(extent)”

等,其语义格性质还不清楚,需要进一步研究。

袁毓林(2002)提出过一个论元角色清单,把客体论元(相当于这里的客事)和与事合并,考虑直接关联的实例:

小邵给汪锋一本书

由于与事和客事可以在同一个动词中直接关联,我们把两者分开。根据同样理由,时间、时段和处所也要分开。把时间和处所分开是大多数学者的主张,也可以得到直接关联原则的支持。我们把时间和时段分开,也是坚持直接关联原则的结果。比如:

他昨天等了我一个小时

“昨天”和“小时”和动词“等”直接关联,需要分开。

直接关联原则会遇到一些例外,这些例外还没有得到合理的解释。现在来观察一些比较典型的例外:

我在学校吃食堂(两个处所)

*我用筷子吃大碗

我在食堂吃大碗

*我用大碗吃食堂

“学校”和“食堂”都是处所,在第一个句子中同时和“吃”直接关联。这和直接关联原则不一致,因为上面说同一个语义格的变体(这里是处所)不应该同时和一个动词直接关联。一个可能的解释是,“吃食堂”可能要理解成“吃食堂的饭”,V语符库中要标注,因为不是所有的V都可以这样用,这是不周遍的。这样的解释似乎又有这样的例外:

米饭我吃食堂

如果“吃食堂”要理解成“吃食堂的饭”，怎么解释上面同一语义格在不同位置上共现？有可能下面两个实例是平行的：

米饭我吃食堂

水果我吃香蕉

再考虑一个例外：

在北方的农村，老百姓经常在家门口晒太阳

这里出现了两个处所 NP，似乎也和直接关联原则不一致。有可能“北方的农村”不是“晒”的处所，而是“老百姓”所在的处所或整个 S 的处所，因此“北方的农村”不能直接出现在“晒”的左边：

* 老百姓经常在北方的农村晒太阳

再来考虑一些实例：

* 我吃饭和筷子

我用筷子吃饭

这双筷子我吃饭

这只鸡我吃头

这顿饭我吃食堂

以下几个有趣的现象：

1. “吃”后面只能有一个 NP，前面可以有两个 NP。

2. 工具、受事、施事可以共现。

3. 关系成分“这只鸡”可以出现,但只能出现一个关系成分,关系成分也是平行周遍的,所以不是黏着语义格。

前面说语义格和名词并不是一一对应的,语义格也可以是谓词或句子:

小邵喜欢[_{vp}抽烟]

小邵知道[_s汪锋要去香港]

汪锋通知大家[_s明天上午打篮球]

另一方面,也不是动词才有语义格。下面的名词也有语义格:

[his^{施事}] destruction of [the building^{受事}]

英语中名词的语义格框架也可以看成是由动词转换过来的,并且有平行周遍关系,即只要一个动词有相应的派生名词,也就有相应的格框架。但是哪些动词有派生名词,没有周遍条件,所以在词库中要标注。

汉语中有很多动词和名词有相同的形式,语义格也是对应的。

我们讨论问题	我们对问题的讨论
我们学习英语	我们对英语的学习
我们了解情况	我们对情况的了解

也有些名词并没有相应的动词,但由于语义和行为有关系,也有语义格。

对俄国的战争

袁毓林(1992)曾经对汉语的名词配价做过有价值的研究,也是属于名词的语义结构关系问题。

8.5 动词的格框架和标注

既然语义格关系是初始的,而且不同的动词的格框架的平行周遍条件还没有找出来,格框架(case frame)就应该在词库中标注。Fillmore(1968)也是主张在词库中标注动词的格框架。Fillmore 认为动词首先可以根据格框架进行分类。如:

```
run [__ A]
remove [__ A, __ O]
.....
```

有些格是必选的,有些格是随意的,都要在格框架中描写清楚。比如 open 的分布:

- | | |
|---------------------------------|--------|
| 40. The door(O) opened | 门打开了 |
| 41. John(A) opened the door | 约翰打开了门 |
| 42. The wind(I) opened the door | 风打开了门 |

因此 open 的格框架是:

```
open: +[__ O(I)(A)]
```

这表示 O 是必选的, I 和 A 是任选的。从这种意义上讲, case frame 和 Chomsky(1965, 2. 4. 4, p124)谈到的集合系统(set-system)确实有联系。

语义格的初始性使动词词典的信息又多了一层。以汉语

“咬”的标注为例,应该有:

NP1(A) — 咬 — NP1(O)

这时“咬”的线性结构包含三层初始概念:

1. NP1 和 NP2 是名词或名词短语
2. 在线性结构中, NP1 在“咬”的前面, NP2 在“咬”的后面。
3. NP1 是施事, NP2 是受事。

一个语符 V 往往可以和很多对象发生语义结构关系:

昨天在食堂他送了我一本书

昨天在食堂他买了两个菜

昨天在食堂他哭了两个小时

语义格可以定义为在语言表达中和语符 V 发生语义结构关系的对象。可以考虑把下面 3 个语符 V 的语义格标注如下:

送 [时间、处所、施事、与事、受事]

买 [时间、处所、施事、受事]

哭 [时间、处所、施事、时量]

这些语符 NP 构成了语符 V 的语义格,但不是所有的语符 NP 都要作为语符库中的语义格,因为对于任何表示行为动作的语符 V 来说,只要表达需要,都可以和时间、地点、时量发生语义结构关系,这一点对所有的语符 V 来说基本是周遍的,因此语符 V 和时间、处所产生语义结构关系可以作为语法手册中的一条原则,于是上面的标注可以简化为:

送 [施事、与事、受事]

买 [施事、受事]

哭 [施事]

这3个语符V在语义格上有差异,这种差异是个别事实,不能通过周遍规则或其他参数预测,语言习得需要记住这些特殊事实,因此语符库中需要标注这些语义格。

认知语言学研究对于格标注有重要意义。认知通常是一种周遍规律,找出这些认知规律,词库中的格标注就可以大大简化,因为通过知识结构可以断定哪些语义格是可推导的,不需要标注。比如上面提到的行为动词,都可以在时间空间中发生,时间和空间并不需要在词库中标注。这方面的研究的关键在于弄清楚可推导语义格和不可推导语义格的界限在什么地方。我们认为这个界限可以通过我们前面讨论的平行周遍原则来界定:凡是通过平行周遍原则可以解释的或可以推导的语义格,都不作为语符库中的语义格标注。于是我们可以把语义格分成两种,自由语义格和黏着语义格或固定语义格。语符库中语符V必须标注的语义格称为黏着语义格,不需要标注的是自由语义格。有些语义格是黏着的还是自由的比较容易区分,有些需要进一步研究。比如,“工具”是不是黏着语义格需要深入研究,在我们还没有弄清楚哪些语符V需要工具格以前,工具格应该在V语符库中标注。一般地说,在没有找到语义格的周遍分布条件以前,语义格都应该看做是黏着语义格存放在语符库中。这个道理和判定规则组合是一致的,即不满足平行周遍条件的组合都应该存放在语符库中。

黏着语义格一般情况下不能省略,自由语义格可以省略,以上的“昨天在食堂他送了我一本书”为例:

例句	省略部分
昨天在食堂外边他送了我一本书	
在食堂外边他送了我一本书	时间(自由语义格)
昨天他送了我一本书	处所(自由语义格)
他送了我一本书	时间、处所(自由语义格)
* 昨天在食堂外边他送了一本书	与事(黏着语义格)
* 昨天在食堂外边他送了我	受事(黏着语义格)

当然,是否省略只是一个参考,不能作为确定黏着语义格的标准,因为在特定的语境条件下,各种省略都可能发生。考虑下面实例:

他睡了
他睡大床

从是否可以省略的角度看,“床”不是黏着语义格。但从平行周遍原则看,我们目前还不能概括语符 V 在什么语义特征条件下后面可以带一个处所语义格,很多和“睡”在语义上很相似的语符 V,后面不可以直接带处所语义格:

他在大床上睡着	他睡大床
他在大床上躺着	他躺大床
他在大床上跳着	* 他跳大床
他在大床上爬着	* 他爬大床
他在大床上站着	* 他站大床
他在大床上蹲着	* 他蹲大床
他在大床上呆着	* 他呆大床
他在大床上倒立着	* 他倒立大床
他在大床上藏着	* 他藏大床

“睡、躺”后面带“床”仍然是一特殊事实,所以“床”等处所格需要作为“睡、躺”的黏着语义格。

有些学者区分域内论元和域外论元,或内部论元(internal argument)和外部论元(external argument),施事被作为域外论元。从平行周遍原则看,如果语符 V 是表示施加一种行为的,都可以有施事,这是平行周遍的。假设我们像处理时间、处所一样把施事放在语法手册中处理,那么在语符库标注中要给出语符 V 的“施加行为”特征:

送 [与事、受事],〈施加行为〉

买 [受事],〈施加行为〉

哭〈施加行为〉

和前面相比较,减少了“施事”的标注,却增加了“施加行为”的标注,从这一点上看,标注“施事”还是标注“施加行为”是等价的。但从整个系统看,“施加行为”特征在解释其他句法现象和语义现象上有很大作用,所以标注“施加行为”特征,删除“施事”符合简单原则。因此,从平行周遍的原则看,施事被作为域外论元来处理是有道理的。更主要的是,“施加行为”是一个认知范畴,一般是可以从动词语义中推导出来的,不需要记忆。比如下面的动词都可以推导出施加行为特征:

送、丢、哭、走、跑、跳、拿、抬、举、吃、喝……

当然,如果从实用的角度对“施事、行为”特征都标注,增加赘余特征,也不是不可以。尤其当我们对“施加行为”的理解有争议的时候。比如,有人可能认为“丢”并不是施加行为。

平行周遍原则只是我们用来判定哪些与动词有关系的搭配成分需要存放在词库中,至于动词和相关成分的语义结构关系

是相当复杂的,同样的句法结构框架,可以有很多不同的语义结构关系。我们来观察动词结构框架的主要类型:

NP—V

[小王^{施事}]来了
[蒋介石^{客事}]死了

NP—V—NP

[小王^{施事}]砸[杯子^{受事}]
[小王^{施事}]认识[小邵^{对象}]
[小王^{施事}]吃[食堂^{处所}]
[小王^{施事}]吃[大碗^{工具}]
[小王^{施事}]炒[菜^{结果}]
[小王^{施事}]拿[第一^{结果?}]

[小王^{施事}]走[小路^{处所}]
[小王^{施事}]走[S形^{方式}]
[小王^{施事}]去[北京^{处所}]
[小王^{施事}]坐[沙发^{处所}]

[他^{客事}]姓[张?]
[小王^{客事}]是[学生?]
[中国^{客事}]属于[亚洲?]
[幻想^{客事}]成为[泡沫^{结果?}]
[刘少奇^{客事}]当[主席?]

NP—V—NP—NP

[小王^{施事}]教[小张^{与事}][汉语^{客事}]
[小王^{施事}]问[小张^{受事?}][问题?]

[小王^{施事}]请教[小张^{受事?}][难题?]

[小王^{施事}]要[小张^{受事?}][一本书?]

[小王^{施事}]告诉[小张^{与事}][一道题?]

[小王^{施事}]借[小张^{与事}][一本书^{受事}]

[张三^{施事}]偷[小张^{与事?}][一本书^{受事}]

[小王^{施事}]花了[小张^{与事}][100元钱^{受事}]

[小王^{施事}]叫[小张^{受事?}][张老师?]

NP—V—V

[小王^{施事}]做[研究^{行为}]

NP—V—VP

[小王^{施事}]喜欢[看小说^{行为}]

[小王^{施事}]开始[穿西装^{行为}]

NP—V—NP—VP

[父亲^{施事/说服}]说服[儿子^{受事/说服 ∩ 施事/学}]学经济

[母亲^{施事/同意}]同意[儿子^{受事/同意 ∩ 施事/学}]学中文

[儿子^{施事/答应 ∩ 施事/学}]答应[父亲^{受事/答应}]学中文

[大家^{施事/选举}]选举[王老五^{受事/选举 ∩ 施事/当}]当村长

NP—V—S

[小王^{施事}]知道[东京在日本^{陈述}]

[小张^{施事}]听说[东京在日本^{陈述}]

[小李^{施事}]正在查阅[东京在哪里^{疑问}]

NP—V—NP—S

[小王^{施事}]告诉[小张^{与事?}][东京在日本^{陈述}]

[小张^{施事}]问[小王^{受事?}][东京在哪里^{疑问}]

NP—PP—V—V

[小王^{施事}]正在[对物理学^{对象}]展开[研究^{行为}]

[小王^{施事}]正在[对财产^{受事}]进行[清理^{行为}]

NP—V—PP

小鸟飞[向远方^{处所}]

小王住[在二楼^{处所}]

以上只是比较常见的情况,其中对题元角色的标注有不少是有分歧的。正因为语义格的标注相当复杂,现在还没有做深入的研究,所以还不能抛弃动词的次范畴化结构框架。以上标注的复杂性从某个侧面说明 Chomsky(1965)提出的通过动词的严格次范畴给出动词的框架,仍然是有用的。

语义格标注还应该包括论元层次的标注。前面说生成语法中还区分内在语义格或外在语义格,即内部论元(internal argument)和外部论元(external argument),或域内论元和域外论元。我们从平行周遍的角度谈到如何在词库中处理固定论元和自由论元,其实这已经显示这两种论元也代表了两种不同的论元层次。考虑下面二价论元结构:

[小王^{施事}]摔了杯子

[小王^{客事}]摔了腿

以上“小王”是施事还是客事,和动词与受事组合后的意义有关系,这再次证明主动句中处在主语位置上的论元相对于动词后

宾语位置上的论元来说是外部论元,这个问题的本质就是论元和动词的关系是有层阶的,即不同的论元和动词的紧密结合程度并不一样。内部论元和外部论元的区分已经显示了题元层阶(thematic hierarchy)的存在,不过论元层阶不仅仅是内部和外部的二分问题,论元层阶可能是多层的。观察实例:

昨天汪锋砸了杯子

砸了杯子

· 汪锋砸了

· 昨天砸了

? 昨天砸了杯子

· 昨天汪锋砸了

从以上实例的合法程度可以看出论元的呈现是有等级的,“杯子”和动词的关系是最内在的,因此是最内层的,如果“杯子”没有出现,“汪锋”或“昨天”直接和动词组合都不合法。这证实了 Wells 早期的一种观点:动词和宾语是谓语的直接成分,主语和谓语是句子的直接成分。后面的 X 杠理论实际上也部分接受了 Wells 的这种观点。不过当时 Wells 并没有区分句法结构的层次和语义结构的层次,或者说没有区分表层结构的层次和深层结构的层次。考虑下面两种层次:

(汪锋)([砸了][杯子])

(杯子)([汪锋][砸了])

在前一个句子中“汪锋”处在外层,“杯子”处在内层,在后一个句子中“杯子”处在外层,“汪锋”处在内层。但在语义结构或论元结构层次上,“杯子”总是内层,“汪锋”总是外层,理由就是呈现等级。

沿着论元结构的角​​度再往外围看,“汪锋砸了杯子”比“昨天砸了杯子”更合法,因此“汪锋”应该比“昨天”在语义结构关系的层阶上更核心一些。于是我们可以根据论元呈现等级进一步给论元标注层阶。比如:

[砸]受事)施事)时间)

以上并不是线性关系,也不是表层的树形图关系,而是和动词联系紧密程度的层阶关系。因为在表层关系上,“杯子”也可以在最外层,时间“昨天”也可以在最内层:

杯子[汪锋[昨天砸了]

也可以考虑用下标的数字表示论元的层阶关系:

[砸]受事)₁,施事)₂,时间)₃

下标 1 表示第 1 层,下标 2 表示第 2 层,下标 3 表示第 3 层。

再考虑更多的论元:

王岳川昨天在会议室给邵永海一幅字

给一幅字

给邵永海

给邵永海一幅字

王岳川给一幅字

王岳川给邵永海一幅字

昨天给邵永海一幅字

王岳川给邵永海

·王岳川给

- 在会议室给
- 昨天给
- 昨天在会议室给
- 王岳川昨天给
- 王岳川在会议室给
- 王岳川昨天在会议室给

这里的论元层阶应该是：

[给]与事∪受事)₁施事)₂时间∪处所)₃

论元层阶并不一定是跟题元角色对应的，比如下面的例子中处所、路径比施事或客事更内层：

- 小王住过大房间
- 小王爬了两座山
- 小王转了两个圈

以上是按照必须呈现的等级来判定论元的层阶。如何一般地判定论元的层阶，需要进一步研究。通常所说的内部论元和外部论元，也还没有找到严格的标准进行判定，主要凭借经验。生成语法的语杠理论(X 杠理论)是从补足语(Comp)和限定语(Spec)的角度来确定域内论元和域外论元，补足语是域内论元，限定域是域外论元，绕开了论元角色，有一定的好处。但如何判定补足语比限定语更内层，也需要研究。后面我们讨论 X 杠理论的时候会展开一些分析。

从以上分析可以看出，论元的层阶和如何在词库中标注论元并不是完全相同的两个问题。在词库中需要标注的论元是不满足周遍条件的论元，论元的层阶则是根据呈现等级来标注。

8.6 语义结构变体标注

具有相同论元的动词,语义结构变体是不同的。汉语基本上可以分出几种类型。一种是可以直接形成线性结构:

孩子砸了杯子
老师了解学生
妈妈给孩子生日礼物

一种必须要通过介词引入论元:

老师马上跟学生见面

还有一种是通过非线性的方式出现:

小邵将汪锋的军
老师见学生的面

前面说过,需要根据直接关联原则确定基本的语义格。但进一步怎么处理动词语义结构的个性,是格语法遇到的问题,因为即使语义结构关系和深层线性结构平行的动词,它们的转换方式或表层句式也可能不同。比较下面的施事受事关系:

小王骂了小张→小张被小王骂了
小王看了小张→‘小张被小王看了

“看”和“骂”后面的语义格相同还是不相同,根据标准的宽严会有不同的结果。无论采取宽或严的态度,一个可以转换成“被”

字句,一个不可以,这种可观察事实是存在的,应该得到充分的描写。在没有找出被字句转换的平行周遍条件以前,哪些 V 语符可以转换成被字句需要在语符库中标注,这就是语符库的语义结构变体标注。

Levin 等(2005)总结了他们所提出的事件结构理论(event structure),认为可以解决格语法遇到的问题。Levin 等的事件结构理论提供了一种不同的研究思路,不过动词语义结构的个性仍然不能绕过。关键问题就在于,不同的语符 V,语义结构框架相同但可以有不同的变体,这是转换的一个重要参数。一个语符有哪些语义结构变体,是语符 V 的特殊属性。比较下面的实例(有的需要一定的上下文才能出现):

N1 + V + N2	N1 + 把 + N2 + V	N2 + 被 + N1 + V	N2 + N1 + V	N1 + V + 的 + N2	V + N2 + 的 + N1
李四砸 了杯子	李四把 杯子砸 了	杯子被 李四砸 了	杯子李 四砸了	李四砸 的杯子	砸杯子的 李四
李四读 了课文	李四把 课文读 了	课文被 李四读 了	课文李 四读了	李四读 的课文	读课文的 李四
李四住 四楼	李四把 四楼住 了	四楼被 李四住 了	四楼李 四住了	李四住 的四楼	住四楼的 李四
李四喝 了茶	李四把 茶喝了	茶被李 四喝了	茶李四 喝了	李四喝 的茶	喝茶的李 四
李四卖 了电影 票	李四把 电影票 卖了	电影票 被李四 卖了	电影票 李四卖 了	李四卖 的电影 票	卖电影票 的李四

续表

N1+V +N2	N1 + 把 + N2+V	N2 + 被 + N1+V	N2 + N1+V	N1+V+ 的+N2	V+N2+的 +N1
李四走 田字		田字被 李四走 了	田字李 四走过	李四走 的田字	走田字的李 四
李四跑 弯道		弯道被 李四跑 了	弯道李 四跑过	李四跑 的弯道	跑弯道的李 四
李四睡 地铺		地铺被 李四睡 了	地铺李 四睡过	李四睡 的地铺	睡地铺的李 四
李四恨 几何			几何李 四恨过	李四恨 的几何	恨几何的李 四
李四爱 动物			动物李 四爱过	李四爱 的动物	爱动物的李 四
李四懂 武术			武术李 四懂了	李四懂 的武术	懂武术的李 四
李四讲 英语			英语李 四讲了	李四讲 的英语	讲英语的李 四
李四开 小车			小车李 四开过	李四开 的小车	开小车的李 四
李四有 杯子			杯子李 四有	李四有 的杯子	有杯子的李 四
李四在 书房			书房李 四在过	李四在 的书房	在书房的李 四
邻居叫 李四				邻居叫 的姓名	叫李四的邻 居

续表

N1+V +N2	N1 + 把 + N2+V	N2 + 被 + N1+V	N2 + N1+V	N1+V+ 的+N2	V+N2+的 +N1
弟弟像 李四				弟弟像 的人	像李四的弟 弟
邻居姓 李				邻居姓 的姓氏	姓李的邻居
李四是 专家					是专家的 李四
李四算 学者					算学者的 李四

以上语义结构变体有一种蕴含关系,如果有“N1+把+N2+V”格式,通常也有其他格式。如果有“N2+被+N1+V”,通常也有“N1+V+的+N2、N2+N1+V”。如果有“N2+N1+V”,通常也有“N1+V+的+N2”。如果有“N1+V+的+N2”,通常也有“V+N2+的+N1”。以上规律反过来不成立,所以蕴含过程可以概括为:

$$N1+把+N2+V \supset N2+被+N1+V \supset N2+N1+V \supset N1+V+的+N2 \supset V+N2+的+N1$$

在语义结构变体之间都有特殊的限制,比如“把”字式和“被”字式的 N2 是有定的,V 的前后有一定的成分,等等,这些条件可以在语法手册中描写。但即使在满足这些条件的情况下,哪些语符 V 有哪些语义结构变体,平行周遍条件还没有完全找到。因此 V 语符库中要标注每一个满足 N1+V+N2 框架的语符 V 的语义结构变体,语符库中要注明哪些语符 V 有“把”字式,哪些有“被”字式,等等。

上表显示所有语符 V 都有“V+N2+的+N1”结构变体,因此这种变体格式是普遍的,可以在语法手册中描写,不必在语符库中标注。通过进一步的分析我们发现,不管是什么种类的语符 V,对于“N+VP”,都可以有“VP+的+N”,因此“VP+的+N”可以作为周遍规则存放在生成句法手册中,不必在 V 语符库中标注。

下面考虑另一种语义格关系:

N1(施事)+V+N2(与事)+N3(受事)

李四送张三一本书

李四教张三英语

李四买张三一幢房子

李四卖张三一幢房子

这里的 N1 都是施事。前面已经谈到,在所有的情况下,“N1+V+N2+N3”都有变体格式“V+N2+N3 的 N1”,因此下面只考虑其他变体格式,基本格式“N1+V+N2+N3”和变体格式“V+N2+N3 的 N1”不再列出来:

N1 + 把 + N3 + V+N2	N3+被 +N1+ V+N2	N1 + V +N3+ 的+N2	N1 + V +N2+ 的+N3	N3 + N1+V +N2	N2 + N1 + V+N3
李四把 这本书 给张三	这本书 被李四 给张三 了	李四给 (他)书 的人	李四给 张三的 书	这本书 李四给 张三了	张三李四 给(他)这 本书
李四把 这本书 送张三	这本书 被李四 送张三 了	李四送 (他)书 的人	李四送 张三的 书	这本书 李四送 张三了	张三李四 送(他)这 本书

续表

N1 + 把 + N3 + V + N2	N3 + 被 + N1 + V + N2	N1 + V + N3 + 的 + N2	N1 + V + N2 + 的 + N3	N3 + N1 + V + N2	N2 + N1 + V + N3
李四把这十块钱还张三	这十块钱被李四还张三了	李四还(他)十块钱的人	李四还张三的十块钱	这十块钱李四还张三	张三李四还(他)这十块钱
! 李四把房子卖张三了	! 房子被李四卖张三了	李四卖(他)房子的人	! 李四卖张三的房子	这幢房子李四卖张三	! 张三李四卖(他)这幢房子
! 李四把这本书赔张三了	! 这本书被李四赔张三了	李四赔(他)这本书的人	李四赔张三的书	这本书李四赔张三	张三李四赔(他)这本书
李四把这瓶奶喂小狗了	这瓶奶被李四喂小狗了	李四喂牛奶的小狗	李四喂小狗的牛奶	这瓶牛奶李四喂小狗	小狗李四喂这瓶牛奶
李四把秘密告诉张三了	这个秘密被李四告诉张三了	李四告诉(他)秘密的人	李四告诉张三的秘密	这个秘密李四告诉张三了	
		李四买他一幢房子的人	李四买张三的房子		

续表

N1 + 把 + N3 + V+N2	N3+被 +N1+ V+N2	N1 + V +N3+ 的+N2	N1 + V +N2+ 的+N3	N3 + N1+V +N2	N2 + N1 + V+N3
		李四偷 他一本 书的人	李四偷 张三的 书		
		李四求 (他)一 件事的人	李四求 张三的 事	这件事 李四求 张三	张三李四 求这件事
		李四教 (他)英 语的人	李四教 张三的 英语	英语李 四教张 三	张三李四 教英语
		李四问 (他)问 题的人	李四问 张三的 问题	这个问 题李四 问张三	张三李四 问这个问 题
		李四花 了(他) 十块钱 的人	? 李四 花父亲 的钱		
		这本书 花了(他) 十块钱 的人			
		李四欠 了(他) 十块钱 的人	李四欠 张三的 十块钱		

不同的 V 往往有不同语义格框架变体,大概是因为词汇意

义不同。比如,“买”有取得的语义特征,“卖”有给予的语义特征,其他都相同,但语义结构变体就很不同。可见每个语符 V 的语义结构框架变体都要作独立描写。

有些语符 V 既有“取得”的语义特征,也有给予的语义特征,比如:

李四借张三这本书

李四租张三一间房子

但转换成不同的语义结构变体以后,语义特征受到限制:

N1 + 把 + N3 + V + N2	N3 + 被 + N1 + V + N2	N1 + V + N3 + 的 + N2	N1 + V + N2 + 的 + N3	N3 + N1 + V + N2	N2 + N1 + V + N3
给予	给予	取得、 给予	取得、 给予	给予	取得、给 予
李四把 这本书 借张三 了	这本书 被李四 借张三 了	李四借 (他)书 的人	李四借 张三的 书	这本书 李四借 张三了	张三李四 借这本书
李四把 这间房 子租张 三了	中间房 子被李 四租张 三了	李四租 (他)房 子的人	李四租 张三的 房子	这间房 子李四 租张三 了	张三李四 借这间房 子

这说明“把”字式、“被”字式和 N3 + N1 + V + N2 格式不具有“取得”的语义特征,所以“N1 + V + N2 + N3”中的语符 V,凡是有“取得”语义特征的,都不能进入“把”字式、“被”字式和 N3 + N1 + V + N2 格式:

N1 + 把 + N3 + V + N2	N3 + 被 + N1 + V + N2	N1 + V + N3 + 的 + N2	N1 + V + N2 + 的 + N3	N3 + N1 + V + N2	N2 + N1 + V + N3
取得	取得	取得	取得	取得	取得
		李四买 过他一 幢房子 的人	李四买 张三的 房子		张三李四 买过他一 幢房子
		李四偷 过他一 本书的 人	李四偷 张三的 书		张三李四 偷过他一 本书
		李四取 过他女 儿的人			张三李四 取过他一 个女儿
		李四收 过他礼 物的人	李四收 张三的 礼物		张三李四 收过他礼 物
		李四拿 过他钱 的人	李四拿 张三的 钱		张三李四 拿过他钱
		李四得 过他好 处的人	李四得 张三的 好处		张三李四 得过他好 处

考虑到这种限制,只要 V 语符库中只标注有“取得”义而没有标注“给予”义的语符 V,不必再标注是否具有“把”字式、“被”字式和 N3 + N1 + V + N2 格式。

“李四问张三一个问题”中的“问”,没有“取得”的语义特征,也不能进入“把”字式、“被”字式,但可以进入 N3 + N1 + V + N2,比如“这个问题李四问过张三”,这说明“把”字式、“被”字式

的语义限制和 $N3+N1+V+N2$ 的语义限制不完全相同。“把”字式、“被”字式中的语符 V 没有“取得”的语义特征,由于“买”类语符 V 和“问”类语符 V 有“取得”的语义特征,所有都不能进入“把”字式和“被”字式。 $N3+N1+V+N2$ 要求没有“取得”的语义特征,“买”类语符 V 有“取得”的语义特征,不能进入这种格式,“问”类语符 V 没有“取得”的语义特征(也没有“给予”的语义特征),可以进入这种格式。

经过以上分析可以看出,对于出现在 $N1+V+N2+N3$ 格式中的“卖、买、问”这类语符 V ,只要对是否有“取得”、“给予”的语义特征进行标注,然后能在语法库中描写出变体规则,就可以预测部分结构变体,这些结构变体不必在 V 语符库中标注。

“李四叫张三大哥”也具有“ $N1+V+N2+N3$ ”的语义结构框架,但“李四叫张三大哥”的语义结构变体和上面实例的结构变体有很大区别。上面结构变体的转换方式是相同的,只是有的语符 V 缺少某些变体,但“李四叫张三大哥”的有些结构变体转换方式比较特殊:

$N1+把$ $+N2+$ $V+N3$	$N2+被$ $+N1+$ $V+N3$	$N1+V$ $+N3+$ 的 $+N2$	$N1+V$ $+N2+$ 的 $+N3$	$N3+N1$ $+V+N2$	$N2+N1$ $+V+N3$
李四把 张三叫 大哥	! 张三 被李四 叫大哥	李四叫 (他)大 哥的人		那个人 李四叫 他大哥	张三李四 叫大哥
李四把 张三称 教授	! 张三 被李四 称教授	李四称 (他)教 授的人		那个人 李四称 (他)教授	张三李四 称教授

转换的特殊性主要在“把”字句和“被”字句。比较:

N1 + V + N2 + N3	N1 + 把 + N3 + V + N2	N3 + 被 + N1 + V + N2
李四送张三一本书	李四把这本书送 张三	这本书被李四送张三 了
N1 + V + N2 + N3	N1 + 把 + N2 + V + N3	N2 + 被 + N1 + V + N3
李四叫张三大哥	李四把张三叫大 哥	! 张三被李四叫大哥

语符“送”是把 N3 作为“把”的支配对象,而语符“叫”是把 N2 作为“把”的支配对象。

8.7 论元结构和次范畴结构的相互独立性

Chomsky(1965)认为深层结构是线性的。前面我们已经论证,深层结构是通过改写规则派生出来的,而改写规则本质上是扩展替换的结果,因此深层结构确实是线性结构。不过,Chomsky(1972)又认为语义格的信息已经在深层结构中,我们前面的分析证明并非如此,语义格是一个初始层面。

现在的问题是,语义结构或论元结构是否可以推导出深层结构或线性结构? Fillmore(1968)认为可以取消深层结构,语义格已经有深层结构的信息。William(1981)给出了一个实现规则(Realization Rule),即有了论元的充分信息,次范畴结构中的语类及其位置都可以推导出来。以英语 give 为例,其域内论元信息为:

Theme:(NP)

Goal:(NP, PP_{to})

Goal:(NP2)

这里含义是,受事 Theme 总是以 NP 的方式出现,与事 Goal 可以以介词 to 的宾语 NP 形式出现,也可以以 NP 的方式出现,但要占据第 2 个名词的位置,即 NP1+V+NP2+NP3 中的 NP2。比如:

Mary gave[John^G] [a book^T]
 Mary gave[a book^T] [to John^G]
 * Mary gave[a book^T] [John^G]

Chomsky(1986a,p86)也有用论元结构代替次范畴结构的意向。实现规则在英语中是否普遍有效还不清楚。至少在汉语中会遇到困难,同样的论元角色,在不同的动词中可以有不同的初始位置:

将[小王^T]的军
 和[老师^T]见面
 拜访[老师^T]

有些动词,同样一个论元角色可以有几种不同的线性顺序,这些线性顺序都可以通过叠加得到,因此也可以看成是初始位置,这些不同的初始位置很难通过论元角色推导出来。比如:

挂[一幅画^T][在墙上^L] (L=Locative)
 挂[在墙上^L][一幅画^T]
 [在墙上^L]挂[一幅画^T]

送[汪锋^G][一本书^T] (G=Goal, T=Theme)
 送[给汪锋^G][一本书^T]
 送[一本书^T][给汪锋^G]

有些位置的论元角色不太好判定,但语类显然是不同的。考虑汉语“知道”的论元结构:

X 知道 Y

Y 的论元角色可能被看成经验的对象,也可能被分成经验的对象和经验的事件,无论采取哪种情况,语类都不同,这是很确定的:

NP1 + 知道 + NP2	张三知道李四
NP1 + 知道 + S	张三知道李四去过上海

“喜欢”的域内论元也可以有不同的语类:

NP1 + 喜欢 + NP2	小王喜欢鱼
NP1 + 喜欢 + VP	小王喜欢吃鱼
NP1 + 喜欢 + VP ∪ NP2	小王喜欢烤鱼(有歧义)

以上分析说明从论元结构不一定能够推导出次范畴框架或语类结构。

由于论元结构框架不能取代语类结构或次范畴框架,又由于深层结构是在次范畴框架基础上通过短语结构规则生成的,本质上是通过初始连接、替换、叠加等线性程序形成的,因此论元结构也不能取代深层结构。深层结构是转换得以展开的基础。是否可以直接在论元结构的基础上开始转换?回答是否定的。从优选角度看,这样的语法全部是转换规则,非常复杂。从必要性上看,由于语法必须包含递归性,因此必须要有初始连接、替换和叠加等线性程序,而这些线性程序的生成结果就是深层结构。考虑“张三会见李四的父亲”这样的句子的生成过程,我们至少需要下面几条规则:

$$S \rightarrow NP + VP$$
$$VP \rightarrow V + NP$$
$$NP \rightarrow NP \text{ 的 } NP$$

这样的规则可以得到：

$$NP + V + NP$$
$$NP + V + (NP \text{ 的 } NP)$$

以上两个表达式的论元关系都是相同的,但深层结构并不相同,因为所使用的短语结构规则不同。深层结构是指通过短语结构规则派生出来的结构,而论元结构并不需要考虑短语结构的派生过程,是词库中标注的抽象信息。在词库中,显然不能包含“NP + V + (NP 的 NP)”这样的信息,否则就等于没有严格区别词库和规则。下面一组句子的论元结构都相同,但深层结构不同:

张三写过书

张三写过一本书

张三写过一本英语书

张三刚写过一本英语书

可见,只有具有扩展替换性质的改写规则有递归性。反过来说,一个语言要有生成性,必须要有递归性的改写规则。而有递归性的改写规则,就会有深层结构。

深层结构是由短语结构规则获得的。短语结构规则本质上就是替换规则或建立在替换规则上的扩展规则。一个语言的句子不可能不用替换规则,因此一个语言不可能不用短语结构规则。改写规则或短语结构规则的存在是深层结构存在的理由,由短语结构规则形成的序列也就是深层结构或线性结构。

在使用改写规则生成“张三会见李四的父亲”以前,不可能先插入词项形成下面结构:

张三会见父亲

因为词项一旦插入,就没有再使用改写规则的机会了。由此可见,必须要有一个进行替换或使用短语结构规则的深层结构。当我们用“NP 的 NP”替换“NP”时,已经假定“NP 的 NP”的存在。是否可以像生成语义学那样,把“张三会见李四的父亲”分成两个原子句:

张三会见父亲

父亲是李四的

然后再引入连接转换把这两个句子连接起来?按照这种生成过程,先得有一个转换:

父亲是李四的→李四的父亲

然后再用“李四的父亲”去替换“张三会见父亲”中的“父亲”。这本质上也是改写规则。

可以说,深层结构是必要的,因为短语结构规则是必要的,替换规则是必要的。后来的 X 阶理论尽管试图减少冗余信息,但并没有从根本上取消短语结构规则,因此也没有取消深层结构。

以上的递归性短语结构规则是否可以用转换规则来描述?如果能,也可以取消深层结构。下面的研究证明不能。Chomsky (1957,1965)是把复杂句子分解成简单句的。Chomsky(1957)的改写规则,都是为简单句用的,复杂句都是通过简单句的连接转换形成的。尤其是 Chomsky(1965,p129)是一个很好的实例:

$\langle_{s_3}$ the man($_{s_2}$ who persuaded John [$_{s_1}$ to be examined by a specialist]) was fired $\rangle$

Chomsky(1957,1965)没有这样的形式:

$NP \rightarrow NP + NP$

$NP \rightarrow NP + PP$

那么名词短语的递归过程是怎么形成的? 是通过小句嵌套。Chomsky(1965,p107)有下面一些改写规则,都是把复杂片断分析成小句:

$VP \rightarrow V \cdot S'$

$NP \rightarrow (Det)N(S')$

当然,Chomsky(1965)并没有把 many、every 这样一些量词分析成谓词,而把它们看成是 determiner。后来有的学者把 many books 这样的片断理解成从 books are many 转换过来的。

下面再从简单句向包孕句扩展的角度来分析生成性的本质。生成语法认为包孕句是从简单句经过转换获得的。下面我们要证明,由简单句向包孕句扩展,必须用短语结构规则,转换规则只起到变形的作用。考虑:

傅刚在吃小邵买的苹果

可以把这个复杂句子看成是下面两个简单句组合的结果:

傅刚在吃苹果

小邵买了苹果

嵌套的结果是：

NP+V+[S+NP]
傅刚在吃[[S]苹果]

具体生成过程可以是：

插入 S	傅刚在吃(小邵买了苹果)苹果
插入“的”	傅刚在吃((小邵买了苹果)的)苹果
删出 S 中的“苹果”	傅刚在吃((小邵买了)的)苹果
删去“了”	傅刚在吃((小邵买)的)苹果

这里有一个生成规则,S 作定语并且 S 中的某个 NP 和 S 所修饰的 NP 同指,S 中的 NP 必须删除。又比如：

NP+V+(S+NP)
老王写字
老张在修桌子
老张在修老王写字的桌子

生成过程如下：

插入 S	老张在修((S)桌子)
	老张在修(老王写字)桌子
插入“的”	老张在修((老王写字)的)桌子
S 中的名词和 S 所修饰的名词没有同指,不删除。	

也可以考虑先对 S 进行名物化转换,但没有插入以前,不知道删除 S 中的哪个名词。一种考虑是不用插入 S 的观点。对于任何

一个 S:

$$\text{NP1} + \text{V} + \text{comp} + \text{NP1}$$

名物化的方式是:

1. 在 S 后加“的”。
2. 把 NP_i 移动到“的”后面, 形成名物化 NP, 其中 NP_i 是 NP 的指称。这个名物化的 NP 可以直接作为单位参加组合。

上述两种方式哪一种好, 涉及评价程序的问题。后一种方法是替换, 前一种方法是插入。为什么可以插入, 显然依据了下面的条件:

$$\text{NP} + \text{V} + (\text{S} + \text{NP}) = \text{NP} + \text{V} + \text{NP}$$

也就是, 必须知道 (S+NP) 和 NP 在上述环境中的分布是等价的, 用 (S+NP) 替换 NP 后并不影响动词的题元关系或次范畴关系。从这种意义上看, 所谓 S 的插入, 本质上是 (S+NP) 替换 NP。

再考虑语义结构的递归性问题:

$$\text{A} + \text{V} + \text{O}$$

是否可以允许 A 和 O 扩展替换? 如果这样, 是否可以取消短语结构规则? 这样的提问是同语反复。因为既然用了扩展替换, 就用了短语结构规则。

到此为止, 我们的分析结果是, 题元结构和建立在次范畴化和线性生成程序基础上的深层结构都是初始概念。不过, 论元

角色的判定问题还没有完全解决,比如:

[_{NP}傅刚^{施事?}]听说[_S明天会出[_{NP}太阳^{施事?}]^{命题}]

和动词“听说、出”有次范畴关系的语类标注的判定标准是比较统一的,但论元角色却存在一定的分歧,所以就目前的研究现状看,如果一定要在初始性上分出等级,次范畴化以及以次范畴化为基础的深层结构是更为初始的概念,因为在论元角色不清楚的情况下,次范畴化框架是进一步确定语法结构关系、语义结构关系的基础。这样说还有一个理由是,次范畴化框架赖以建立的基础是语类,语类是可以在词典中标注的,不仅动词可以在词典中标注语类,和动词发生直接关系的卫星成分也可以在词典中标注语类,但这些卫星成分的语法角色(主语、宾语等)和语义角色(施事、受事等)都不能独立标注,只能根据不同的动词来标注语类。观察下面的实例:

主语 [_N猫^{施事}] 追 老鼠

主语 [_N猫^{受事}] 死了

[狗] 追 宾语 [_N猫^{受事}]

给 间接宾语 [_N猫^{与事}] 一条鱼

“猫”的语类不变,语法角色、语义角色却在变化。因此,我们可以在词库中标注“猫”是什么语类,却无法在词库中独立标注“猫”的语法角色和语义角色。

其实 Chomsky(1965)次范畴化所涉及的动词和卫星成分的关系比 Fillmore(1968)格语法更多。格语法主要考虑动词和名词、命题两种卫星成分的关系,介词被看成引入名词语义格的标记,而 Chomsky 次范畴化所涉及的动词卫星成分种类更多:

次范畴框架	例子
V	elapse
V NP	bring the book
V NP that-S	persuade John that there was no hope
V Prep-Phrase	decide on a new course of action
V Prep-Phrase Prep-Phrase	argue with John about the plan
V Adj	grow sad
V like Predicate-Nominal	feel like a new man
V NP Prep-Phrase	save the book for John
V NP Prep-Phrase Prep-Phrase	trade the bicycle to John for a tennis racket

如何统一解释动词和周围这些卫星成分的语法结构关系和语义结构关系,需要进一步研究,但动词的次范畴框架是研究这些关系的出发点。

8.8 文本中语义格的识别

支配与约束理论由于要在表层结构解释语义,所以空语类的研究相当活跃。空语类是指没有在格位置上出现显性形式的语义格。通过空语类研究,句法的很多属性得到了认识。但有一个重要的方面现在还缺少研究,这就是现实文本中语义格识别的一般原则。

如果我们通过分析得到了语符 V 的次范畴信息,并将这些信息存入词库,我们需要利用这种 V 信息来判定每一个句子中的语义格。Chomsky(1981, p42)讲到过确定题元的两个因素,词典属性和语法功能(GF)。在词库中如何确定语义格,我们在

前面做了初步分析。语法功能就是句子成分。如何判定句子成分,并不是一件容易的工作。尤其是像汉语这样的语言,句子成分缺少标记,或者标记比较隐蔽,判定句子成分更困难。

我们认为词典中的语义结构框架决定了题元或语义成分的基本项目。但是在现实句子中,由于语义结构成分不只是一种位置,就存在确定语义结构成分的问题。可以有下面的基本手段:

1. 有标记的语法形式
2. 语义特征
3. 上下文条件

考虑下面的实例:

[果汁儿_i]我喝了[t_i]

按照 GB 理论(支配与约束理论),“喝”的宾语移动到句首,产生了语迹[t_i],由于句首的“果汁儿”和语迹 t_i同下标,所以是同指,题元或语义格性质一样。问题在于,我们听到的或读到的这个句子的表层形式并没有语迹的标记,我们如何知道是移动的结果,如何知道是从哪里移动过来的,如何知道“果汁儿”是深层结构的宾语或语义结构中的受事?考虑下面的例子:

[苹果]汪锋削了皮儿

[客厅]小邵堆了书

[语言所]我们送了十本杂志

如何找出方括号中成分的语迹,目前依靠的主要是经验,还不是原则。因此,文本中(包括口语文本和文字文本)语义格的识别

是一个有待研究的重要问题。以上“喝”的例句可以通过语义特征来识别,因为在词库中,“喝”的语义结构应该是这样标记的:

施事(NP,动物)+喝+受事(NP,液体)

由于“我”有动物特征,“果汁儿”有液体特征,因此以上句子“果汁儿”理解成受事,“我”识别成施事。这一识别过程不存在问题。

在很多情况下,只靠语义特征无法识别语义格。比较:

喝果汁儿的狗

追孩子的狗

这里的语类序列层次有两个:

a [V][NP 的 NP]

b [VNP 的][NP]

由于“喝”的施事有动物特征,“喝”的受事有液体特征,因此在“喝果汁儿的狗”中,果汁儿和狗的语义格都只有唯一的选择,因此层次也只有 b 的可能。在词库中,“追”的语义结构标注如下:

施事(NP,动物)+追+受事(NP,动物)

由于“狗”和“孩子”都有动物特征,所以施事、受事的配对就有两种情况,因此层次也可以有 a、b 两种,语义格也有两种理解:

[_v 追][_{NP} 孩子的狗^{受事}]

[_{NP}追孩子的][_N狗^{施事}]

这类大家都熟悉的歧义只能通过上下文消除。这里讨论的本质是语义格识别的问题。对于这类动词,由于不同的语义格有相同的语义特征,即使层次相同,语义格有不同理解的情况也是大量存在的:

这只狗孩子也追

(这只狗)([孩子][也追])

(这只狗^{受事})([孩子^{施事}][也追])

(这只狗^{施事})([孩子^{受事}][也追])

这个孩子狗也追

(这个孩子)([狗][也追])

(这个孩子^{受事})([狗^{施事}][也追])

(这个孩子^{施事})([狗^{受事}][也追])

对于“追”这类动词,研究语义格标记是在文本中识别语义格的重要工作。在没有上下文的情况下,语义结构标记是识别语义格的凭借:

[这只狗^{受事}]被[孩子^{施事}]追

[这个孩子^{受事}]被[狗^{施事}]追

介词标记在这里对识别施事和受事有关键作用。根据句法标记判定语义格有强制性。比较

我们被张三耽误了

我们被下雨天耽误了

这里的句法标记“被”有强制作用,即使“下雨天”是时间 NP,也被强化成“施事”或和“张三”相似的语义格。

在没有标记或标记不出现的情况下,需要根据语义特征判定语义格:

过了四级的学生不上英语课

下雨天不上英语课

“过了四级的学生”是一个带有生命语义特征的 NP,可以理解成“上英语课”的施事,“下雨天”由于是时间 NP,只能理解成“上英语课”的时间。

在语义特征可以理解成不同的语义格的情况下,只能根据上下文所提供的条件来判定语义格的角色:

鸡已经吃了

8.9 多项式的判定问题

考虑实例:

我找小王弹吉他

小王的论元角色是什么?这是语义要解释的基本问题。转换生成语法的 GB 理论中的支配与控制理论,本质上就是在研究这里的论元角色问题。比较:

We would like you to stay

We would like [you to stay]

We_i would like to stay

We_i would like [PRO_i to stay] (PRO 并不是 you, 而是 we)

PRO 并不是由移位产生的空语类, 但汉语的语义格标记不丰富, 移位又常常不留形式标记, 因此 GB 理论对汉语多项式的语义格判定有多大作用需要深入研究。下面我们从另一个角度考虑问题。我们先引入一个完备条件:

当一个句子有多个语符 V 时, 这些 V 可以共用某些成分作自己的论元, 除非有上下文, 每个 V 的所有黏着论元都要得到满足。

从 GB 理论的投射原则看, 也应该有这个完备性条件。郭锐 (1995) 提到的合价问题, 也跟完备性条件有关系。因此, 多项式论元的识别本质上是在满足完备性条件下怎样找出不同的黏着论元。先观察下面的实例, 斜杠/后面的动词表示语义格所属动词:

[我^{施事/找}]找[张三^{受事/找, 施事/弹}]弹[吉他^{受事/弹}]

这一实例满足多元语符 V 语义格完备条件, 因为每一项语符 V 的黏着语义格都得到了满足:

找:[我^{施事}],[张三^{受事}]

弹:[张三^{施事}],[吉他^{受事}]

“张三”既作“找”的受事, 也作“弹”的施事, 这就是通常所说的汉语中的兼语式。考虑下面实例:

我找个阴凉处弹吉他

这时的“阴凉处”不是“弹”的施事。这时需要继续往左边寻找施事,最后找到“我”作施事,

[我^{施事}找^{找,施事}弹]^弹找[个阴凉处^{受事}找^{找,处所}弹]^弹弹[吉他^{受事}弹]

于是每一项语符 V 的黏着语义格都得到了满足:

找:[我^{施事}],[阴凉处^{受事}]

弹:[张三^{施事}],[吉他^{受事}]

无论有多少语符 V 并联,对语义格完备条件的要求是一样的。考虑一个有 4 个语符 V 的情况:

李四托我找张三去音乐厅弹吉他

其语义格的标注是:

[李四^{施事/托}]^托[我^{受事/托,施事/找}]^找[张三^{受事/找,施事/去,弹}]^去[音乐厅^{处所}去]^去弹[吉他^{受事/弹}]

这里每一项语符 V 的黏着语义格都得到了满足:

托:[李四^{施事}],[我^{受事}]

找:[我^{施事}],[张三^{受事}]

去:[张三^{施事}],[音乐厅^{处所}]

弹:[张三^{施事}],[吉他^{受事}](由于“音乐厅”不是“弹”的黏着语义格,可以不标注)

有时候句子内部的黏着语义格是找不全的,比如:

明天找个阴凉处弹吉他

最终找出的语义格是:

[明天^{时间}找、弹]找[一个阴凉处^{受事}找,处所^弹]弹[吉他^{受事}弹]

“找”、“弹”的黏着语义格施事都没有得到满足。在这种情况下,必须假定上下文提供了两个语符 V 的施事。再考虑:

下午看望的人我都认识

从语感上可以标注为:

[下午^{时间}看望]看望的[人^{受事或施事}看望,受事^{认识}][我^{施事}认识]都认识
 看望:[人^{施事或受事}]
 认识:[我^{施事}],[人^{受事}]

通过语义格完备条件可以更好的解释一些语义指向问题。考虑:

砍疼了/砍累了/砍断了

这三个格式前面都隐含了 NP,这个 NP 在这三个格式中的语义格是不一样的:

[NP]^{施事或受事或工具}砍,当事^疼砍疼了
 [手]^{受事}/砍,当事^疼砍疼了(砍手,手疼了)
 [手]^{工具}/砍,当事^疼砍疼了(用手砍,砍疼了)

[汪锋]_{施事}砍_{砍, 当事}疼_疼砍疼了

[NP]_{施事或工具}砍_{砍, 当事}累_累砍累了

[汪锋]_{施事}砍_{砍, 当事}累_累砍累了

[手]_{工具}砍_{砍, 当事}累_累砍累了

[NP]_{受事或工具}砍_{砍, 当事}断_断砍断了

[手]_{受事}砍_{砍, 当事}断_断砍断了

[刀]_{工具}砍_{砍, 受事}断_断砍断了

或者从另一个角度说,语义指向识别问题是和论元识别有关系的(陆俭明,1997)。

完备条件是讲必要条件。当两个或两个以上语符 V 出现时,怎样确定哪个 NP 是哪个 V 的施事,哪个 NP 是哪个 V 的受事,是一项重要的工作。目前这方面的工作才刚刚开始。我们下面尝试分析出一个基本规则。

前面分析已经看出,下面的语义结构框架需要在 V 语符库标注:

我知道张三会来	N0+V1+S
我同意张三来	N0+V1+N1+V
.....	

这些情况都属于上面讨论的单项 V 的判定问题。因此多项 V 的语义格判定问题主要集中在下面格式中(以两个 V 为例):

N0+V1+N1+V2+N2

例如:

我找张三弹吉他

张三弹吉他消遣
我去图书馆借书
我借书复印
我拿美元买书
我买书送人
.....

这些 VP+VP 的组合从理论上说是无限的,不可能通过语符库描写,对这类格式中语义格的判定是我们下面要讨论的重点。首先我们引入一个语义结构域的概念。考虑下面两个 V 的情况:

$N0 + V1 + N1 + V2 + N2$

V1 的语义结构域是:[N0+V1+N1]
V2 的语义结构域是:[N1+V2+N2]

比如:

$N0 + V1 + N1 + V2 + N2$	V1 语义结构域: $N0 + V1 + N1$	V1 语义结构域: $N1 + V2 + N2$
我找张三弹吉他	我找张三	张三弹吉他
我找个阴凉处弹吉他	我找个阴凉处	一个阴凉处弹吉他
我找吉他弹	我找吉他	我弹吉他

根据实例分析,可以概括出一条 NP 唯一呈现原则:

当一个 NP 作多项 V 的语义格时,这个 NP 第一次出现以

后,不再重复出现,最多只能以指代语符的形式出现

例如:

[张三_{施事}找,与事_借]找李四借书

[张三_{施事}找,与事_借]找李四借书给[他_{与事}借,给]

*[张三_{施事}找,与事_借]找李四借书给[张三_{与事}借,给]

第1句“借”的与事就是“找”的施事“张三”,不出现,符合独立呈现原则。第2句“借、给”的与事也是“找”的施事“张三”,以指代成分出现,符合独立呈现原则。第3句“借、给”的与事也是“找”的施事“张三”,在句子中重复出现,不符合唯一呈现原则,除非后一个“张三”是指另一个人。容易看出,当指代成分是“自己、他自己”这样一些形式的时候,就和GB理论所讨论的约束理论相似。

从句子生成的角度看,唯一呈现原则的本质是,多项V组合时,和前面NP指称相同的NP要删除或用代词指代,也就是说存在下面的转换过程:

小王找阴凉处+小王在阴凉处弹吉他

→小王找阴凉处弹吉他

小王说服小李 VP+小李报考数学

→小王说服小李报考数学

小王说 S+小王要去上海

→小王说他要去上海

其中斜体的NP和前面NP指称相同,或者被删除,或者被代

词替代。如果强调两个论元的指称不同,则需要都呈现出来:

共享当事	不共享当事
[张三 ^{施事/累、垮}]累垮了	[张三 ^{施事/累}]累垮了[身子 ^{施事/垮}]
[张三 ^{施事/累、病}]累病了	[张三 ^{施事/累}]累病了[身子 ^{施事/病}]
[张三 ^{施事/累、晕}]累晕了	[张三 ^{施事/累、晕}]累晕了[头 ^{施事/晕}]
[张三 ^{施事/累、坏}]累坏了	[张三 ^{施事/累、坏}]累坏了[身子 ^{施事/坏}]
[张三 ^{施事/热、垮}]热垮了	[张三 ^{施事/热、垮}]热垮了[身子 ^{施事/垮}]
[张三 ^{施事/热、病}]热病了	[张三 ^{施事/热}]热病了[身子 ^{施事/病}]
[张三 ^{施事/热、晕}]热晕了	[张三 ^{施事/热}]热晕了[头 ^{施事/晕}]
[张三 ^{施事/热、坏}]热坏了	[张三 ^{施事/热}]热坏了[身子 ^{施事/坏}]

思考问题

1. 如何判定内部论元与外部论元?
2. 如何在多个谓词的组合中判定题元角色?

9. 语义结构的再次非线性化:语义表达式

9.1 深层结构在语义解释上遇到的问题

Chomsky(1965)的标准理论出现后,生成语法在语义解释上出现了不同的模型。Fillmore 格语法是有代表性的一种。就在 Fillmore(1966,1968)研究格语法的同时,其他一些语法学家发展出了生成语义学(generative semantics)。有代表的人物有 P. Postal, J. D. McCawley, G. Lakoff, J. R. Ross。Postal 早些时候和 J. J. Katz 提出了 Katz - Postal 假说,该假说认为深层结构决定意义,此后他们着重研究关系语法(relational grammar)。McCawley 后来研究形式逻辑和自然语言的关系。G. Lakoff 和 J. R. Ross 继续从事生成语义学。Lakoff 和 Ross (1967)在 *Is deep structure necessary* 中将谓词演算和深层结构相类比,并把谓词演算作为语义生成性的起点,认为表层结构和语义解释之间没有必要设立中间结构,语义本身就有生成性。

Chomsky(1965)标准理论中提出的深层结构是为了解释语义,其中一个基本原则是,语义相同的句子,深层结构应该相同,语义不同的句子,深层结构应该不同。由于 Chomsky 的深层结构是由短语结构规则和词汇插入规则生成的,因此只要使用的短语结构规则不同(结构不同),或插入的词项不同,深层结构都应该不同。另一方面,只要使用的短语结构规则相同(结构相同),并且插入的词项相同,深层结构都应该相同。

但实际情况比较复杂。先看深层结构相同语义不同的情况。下面的两个句子,结构和词项都不同,即深层结构不同,但有人认为语义是基本相同的:

Max reminds me of Peter

Max strikes me as similar to Peter

而下面两个句子,基本结构和对应的语类都是相同的,即都使用了相同的短语结构规则,不过由于词项不同,也应该看成深层结构不同,但语义也基本是相同的。

NP + V + NP + PP

John gave a book to Bill

Bill received a book from John

从前面对 Chomsky 的深层结构和 Fillmore 的语义结构看,同义句应该有这样一个前提:

同义句必须在树形结构、语类和词项上都相同。

但事实上,上述两组基本语义相同的句子,第 1 组树形图结构、语类和词项不同,第 2 组树形图和语类同而词项不同。生成语义学要考虑的问题是:有没有一个更深层次的解释来说明这种同义性?

再看深层结构相同但语义不同的情况。下面具有相同的深层结构的句子,意义仍然不同(Lakoff, p238):

(1) Many men read few books.

(2) Few books are read by many men.

这种不同不仅仅是主动和被动不同,还有量词辖域的不同。主动和被动的不同在 Chomsky(1965)的理论中是在深层结构中标注的,量词的辖域却没有在深层结构中标注。除了量词,否定词、连词这类算子(operator)也都存在类似的问题。

Chomsky(1965)的标准理论通过基础部分产生深层结构,是维持线性关系的基础。不过 Chomsky(1965)在概括语法的基本结构时,还提到过一个概念,即语义表达式(semantic representation),Chomsky 认为深层结构通过语义部分(semamtic component)得到语义表达式。至于这个语义表达式是什么,并没有解释和限定。从 Chomsky 的行文看,语义表达式只是对深层结构的一种语义解释,并不是--种语言结构,转换也不是在语义表达式上展开的。生成语义学抓住了语义表达式,从语义表达式直接进行词汇插入规则,并展开转换规则的操作。这样就把语义表达式看成了一种结构,一种可以进行转换操作的结构,这样就否定了 Chomsky 的深层结构。前面我们说过,深层结构是通过连接和替换再插入词汇而生成的线性序列,如果否定了深层结构,就否定了线性关系。从这种意义上说,我们认为生成语义学是语义结构的非线性化。

9.2 语义表达式及其深度

生成语义学首先试图解决不同的深层结构所遇到的同义问题,办法是通过原子命题、词项分解和谓词演算的逻辑结构来描写语义表达式。每一个原子命题都被描写成一个最简单的谓词演算结构。句子基底结构(underlying structure)是语义表达式,语义表达式是原子命题构成的树形图结构(trees 或 phrase - markers)。更明确地说,一个句子的语义解释由下面两个部分组成:

1. 原子命题集合
2. 由原子命题构成的语义表达式

考虑 Lakoff(1967)的例子:

1. Every student expects every student to die
2. Every student expects to die

1 的原子命题有：

- a. student is every
- b. student expects S (S = student will die)
- c. student will be every
- d. student will die

通过下面的树形图(见图 9. 2. 1)可以看出生成语义学的实质：

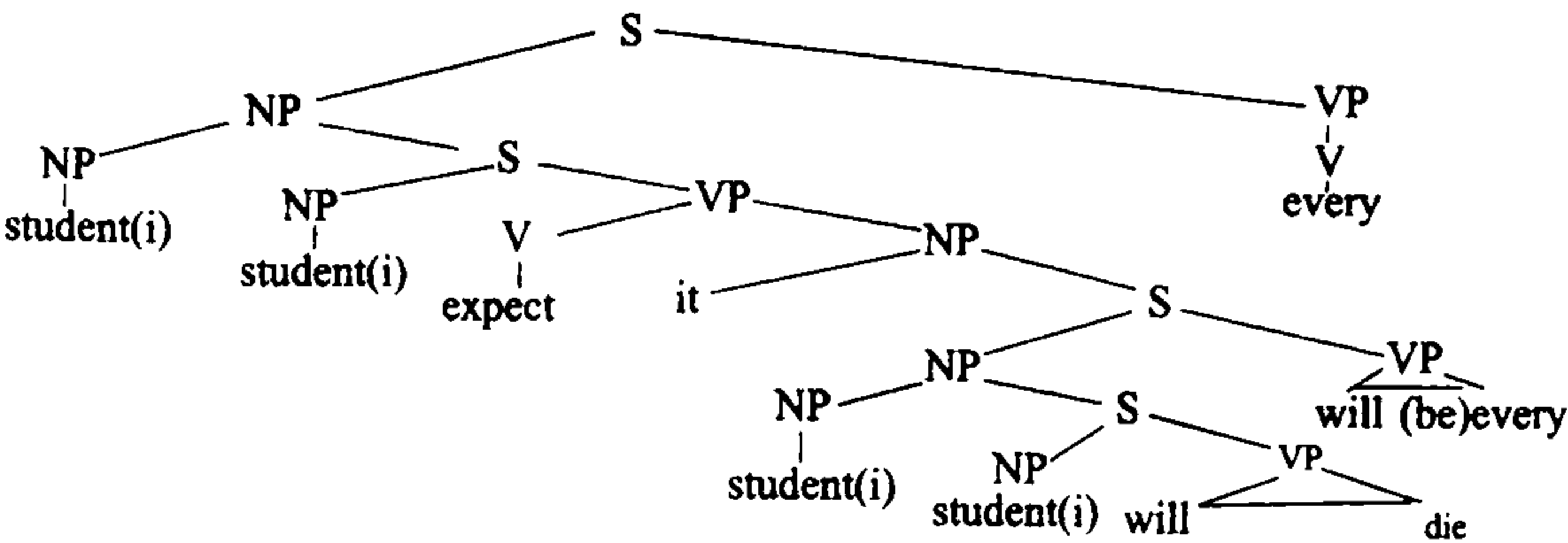


图 9. 2. 1 树形图

再比较下面的图形(见图 9. 2. 2)：

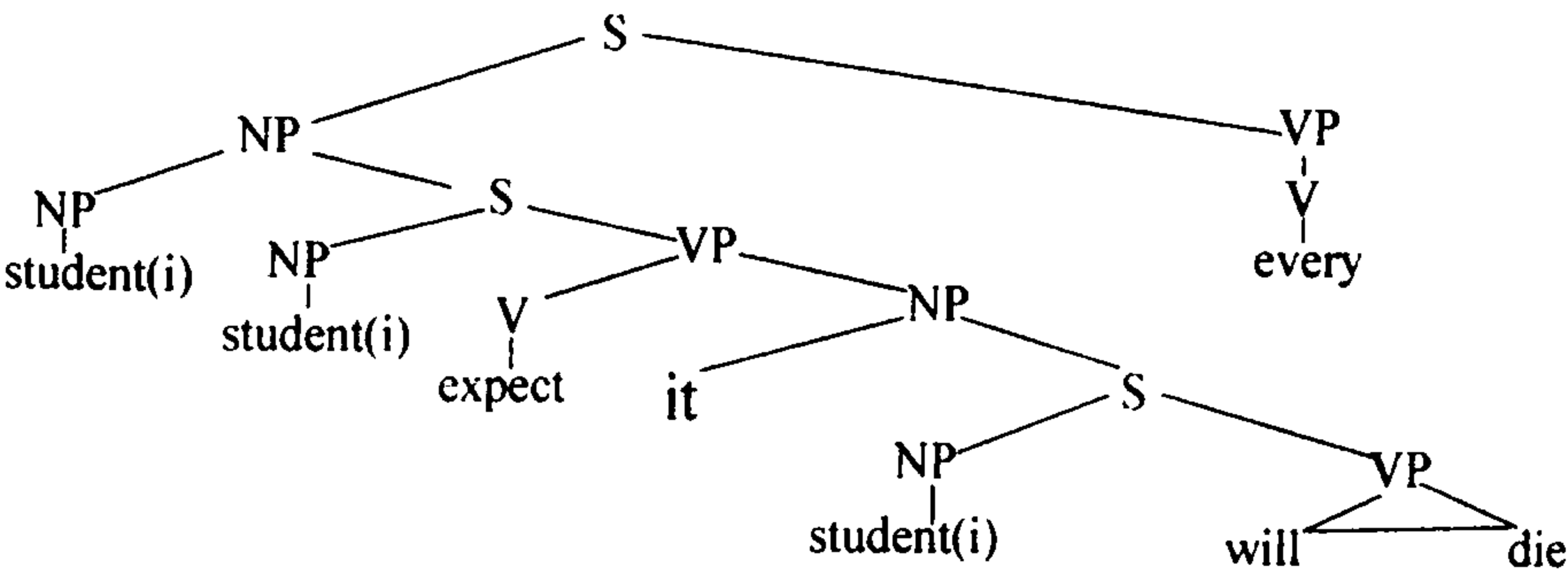


图 9. 2. 2 树形图

概括地说,为了解决不同深层结构的同义性问题和同一深层结构的不同义问题,生成语义学把语义的深度推进了两个

层次:

1. 把修饰语分解成原子命题
2. 把单一动词分解成原子命题

比如, every student 被表达成原子命题 student is every。至于词项的分解, 下面还会谈到。现在结合这两个问题来讨论生成语义学关于语义解释的深化问题。

9.2.1 表达深度 1: 原子命题的组合差别

原子命题是要解决的问题是相同深层结构意义为什么可以不同。再来看前面提到的两个例子(Lakoff, 1972):

- (1) Many men read few books.
- (2) Few books are read by many men.

从 Chomsky(1965)的深层结构和 Fillmore(1968)语义格理论看, 这两个句子的深层结构和语义格关系相同, 因此语义解释都应该相同, 但 Lakoff 认为这是两个意义不同的句子, 基底的语义并不相同。Chomsky(1965)实际上也注意到这是两个不同意义的句子。

按照 Lakoff 的解释, (1)和(2)的语义区别是 many 控制 few (many commands few)和 few 控制 many(few commands many)的区别。Lakoff 用新的短语标记或树形图 (phrase-marker)进行了描写, 这种短语标记的特点之一是把定语分解成原子命题。这样一个句子基底的深度已经比格语法深了, 因为格语法不把定语分解成原子命题。可以用下面的实例说明 Lakoff 的基本分析思想:

many men left → the men who left were many

经过这种分析, men 就和 many 相互 command(控制)了, 这是下面理解 command 的关键。

Lakoff(1972)的另一个基本思路是要把 quantifiers(量词)看成是处在更高层面的成分:

The main point at issue is whether quantifiers in underlying syntactic representations are in a higher clause than the NPs they quantify (as in predicate calculus (谓词演算)) or whether they are part of the NPs they quantify (as they are in surface structure).

这种说法是针对早期的深层结构决定意义说的。Chomsky (1965)把被动和主动看成语义不同, 是在深层结构上加上标记, 以后的转换就不要再改变语义, 所以当时主动和被动在语义上的解释是主动标记和被动标记的区别, 其他不再有语义区别。Lakoff 这里的例子是要说明, 即使在深层结构中标注了主动和被动的区别, 在后来的转换中, 还会引起语义的改变。现在的问题是, 这种改变是否仍然可以在深层结构中标注呢? Chomsky 的标准理论和 Fillmore 的格语法并没有解释这里的区别。按照 Lakoff 生成语义学, 上面的(1)被描写成:

(3) Many are the men who read few books. There are many men who read few books.

其语义结构包括以下原子命题:

a. men(i) are many

- b. books(j) are few
- c. men(i) read books(j)

(1)的树形图或 phrase-markers(见图 9.2.1.1)如下：

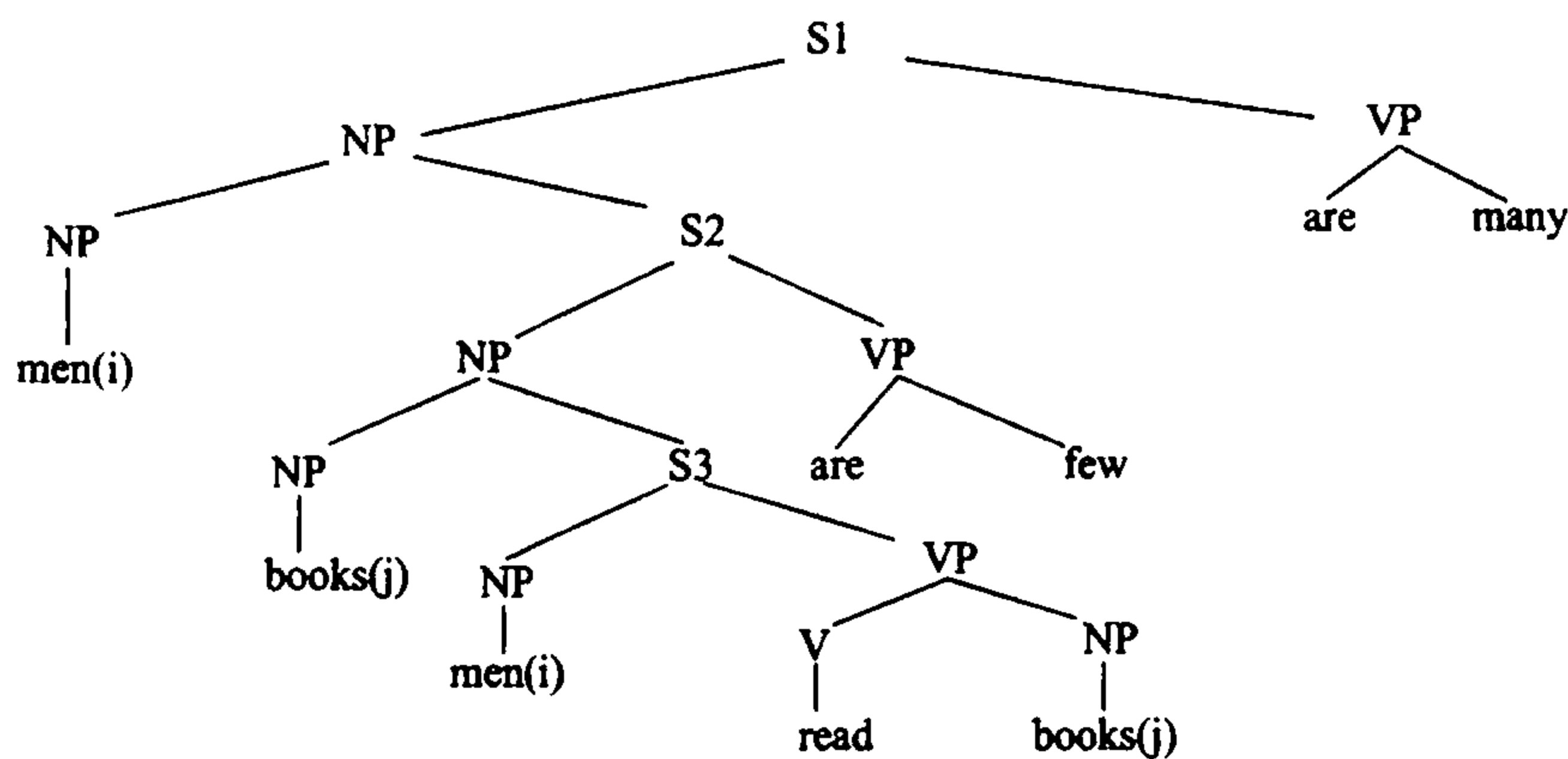


图 9.2.1.1 (1)的树形图

在这个语义表达式中,树形图(tree 或 phrase-marker)已经和短句结构或直接成分不对应。

仅仅把一个句子分解成原子命题是不够的,因为还不能反映出(1)和(2)的差别,这就需要在语义表达式上反映出这种差别。Lakoff(1972,2)主要在做这方面的工作。(1)和(2)有相同原子命题,但(1)的语义表达式和(2)的语义表达式不同。为了区别这两种表达式,以上例句(2)被描写成：

(4) Few are the books that many men read. There are few books that many men read.

其语义结构包括以下原子命题：

- a. books(j) are few
- b. men(i) are many
- c. men(i) read books(j)

(2)的树形图或 phrase-markers(见图 9.2.1.2)如下：

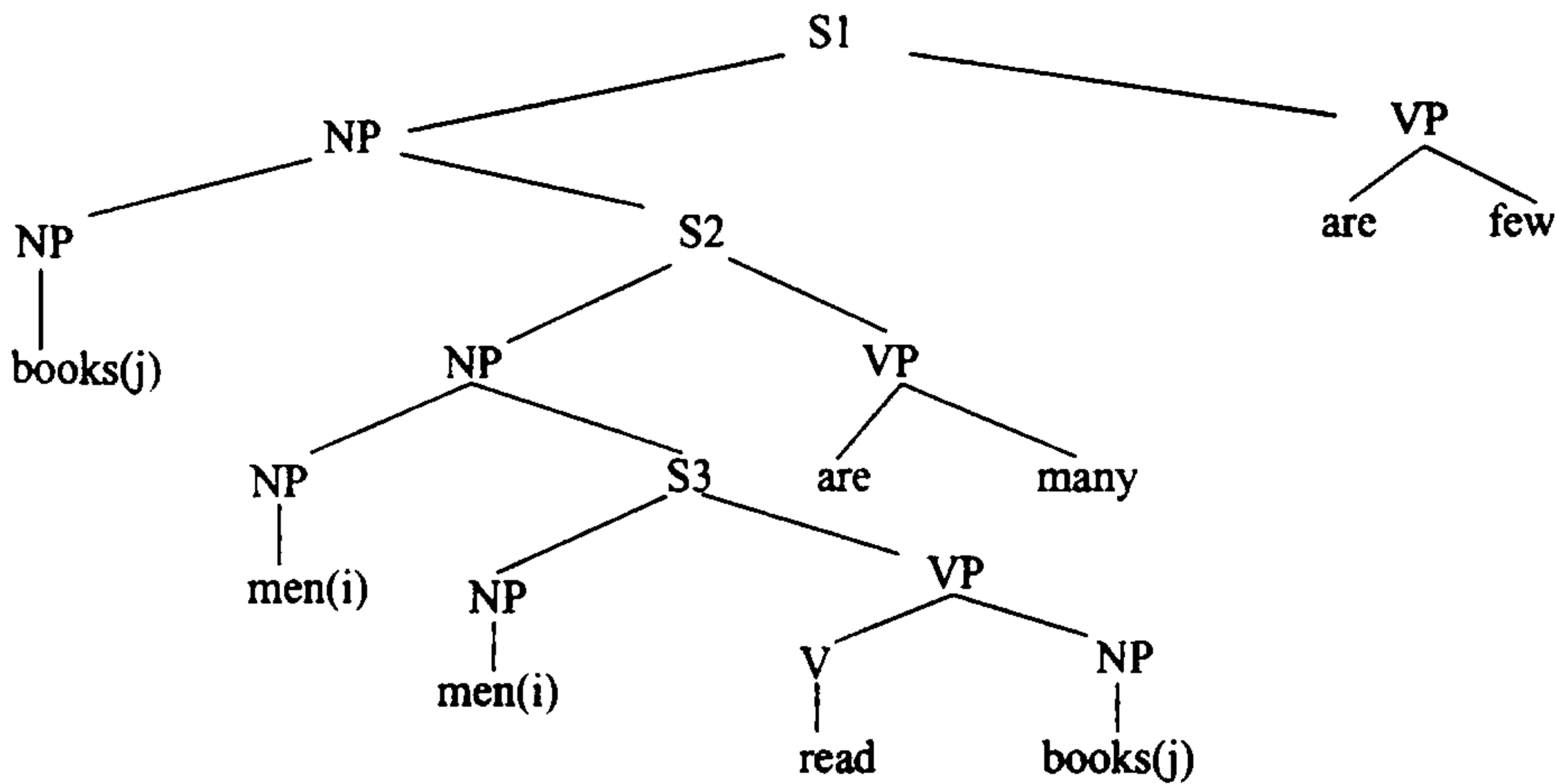


图 9.2.1.2 (2)的树形图

(1)和(2)两个句子的语义结构中的原子命题是相同的,但语义表达式(tree 或 phrase-marker)不一样。或者换句话说,如果两个句子的意义不同,语义表达式(tree 或 phrase-marker)也要被描写成不同。从以上分析可以看出,Lakoff 的 phrase-marker 和建立在 Chomsky(1965)标准理论中的 phrase-marker 并不同。尽管 Lakoff 也用 phrase-marker 这个术语,但都是指他的语义表达式,而 Chomsky 的 phrase-marker 指深层结构及其各个转换阶段的树形图。

这就是说,两个句子尽管有相同的原子命题,由于语义表达式不同,因此语义也不一样。同义句必须有相同的语义表达式,不同义的句子应该有不同的语义表达式。

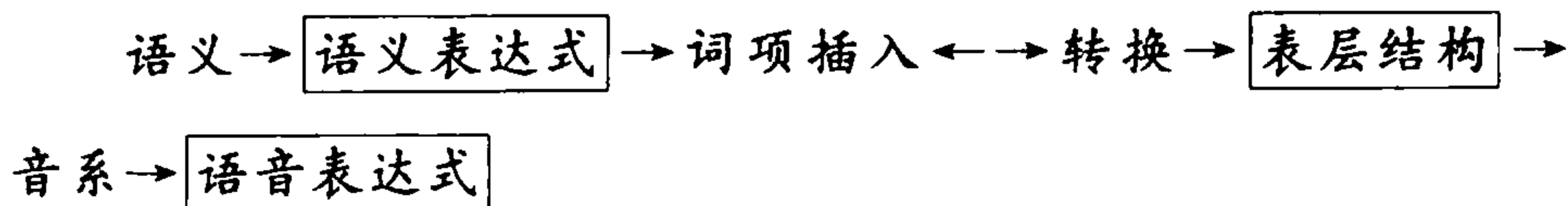
Lakoff 的以上限制是高度抽象的,我把 Lakoff 的做法概括成如下三个方面：

- 1. (1)和(2)是有区别的,这是 surface structure 的区别。
- 2. Chomsky 的深层结构不能表达这种区别。
- 3. 需要通过语义表达式 semantic representation 来区别。

从 Lakoff 的以上陈述可以看出,他是从表层入手在底层解释语义。他的语义表达式是不可观察的,和表层结构不同。分析 Lakoff 的具体研究思路,Lakoff 把 transformation 看成 local derivational constraint。既然量词在 surface structure 上的区别要解释成 underlying structure 的区别,就要有转换或派生限制把底层的不同语义在 surface structure 中体现出来。

Lakoff 要取消 deep structure,自己又用了个 underlying structure,他认为这个 underlying structure 就是 P_1 ,是第 1 个 phrase marker(Lakoff, 1972),即最底层的语义表达式,所以 Lakoff 仍然承认深层的结构,在这一点上和 Chomsky(1965)标准理论是一致的,跟 Katz 和 Postal(1964)的理论框架也是一致的。只是他的 underlying structure 和 Chomsky 的 deep structure 的性质和深度不一样。Chomsky(1965)解释语义的深层结构是用改写规则和词项插入规则决定的,phrase-marker 也是由此决定的。而在 Lakoff(1972)的表达式中,类似 Chomsky 深层结构的东西被进一步分解成语义表达式。

由于没有深层结构,生成的过程必须直接从语义表达式开始。生成语义学结构框架可以概括如下:



$\longleftrightarrow$ 是指,词项插入和转换是相互进行的。这和 Chomsky(1965)标准理论有区别,在标准理论中,词项插入在转换之前完成,这个插入词项的结构就是深层结构。在上面的生成语义学框架中,由于转换以后也可以有词项插入,所以深层结构不存在。

9.2.2 表达深度 2: 词汇分解

生成语义学还提出了词汇分解 (lexical decomposition) 理论来解决不同的深层结构意义可以相同的问题。根据乔姆斯基的标准理论, 符合下面条件的句子, 深层结构不同:

1. 范畴(语类)不同, 词项不同。
2. 范畴(语类)相同, 词项不同。

但是这两种情况下, 语义都可以相同。前面我们已经给出过例子。生成语义学希望解释这样的同义句。尽管按照 Chomsky 的标准理论, 这两组实例都属于深层结构不同(因为语类不同, 词汇插入项不同), 生成语义学认为仍然应该有相同的深层结构或基底结构, 应该有相同的语义结构。所谓抛弃深层结构的动机就在于, 深层结构被词项或语类限定了, 无法说明同义句。可以看出, 生成语义学有一个最为根本的假说, 语义相同的句子, 基底结构都应该相同。

既然词汇项不同仍然是同义句, 要解释这个不同的根源, 就只好超越词汇项(深层结构赖以存在的条件之一)来找同义的深层根据, 这就不得不导致对词项的分解。下面 4 个句子, 从深层结构和格语法的角度看, 深层结构不相同, 语义格也不相同, McCawley (1968) 认为语义表达式是相同的:

- a. John killed Bill
- b. John caused Bill to die
- c. John caused Bill to become dead
- d. John caused Bill to become not alive

语义结构在这里被描写成下图(图 9.2.2.1):

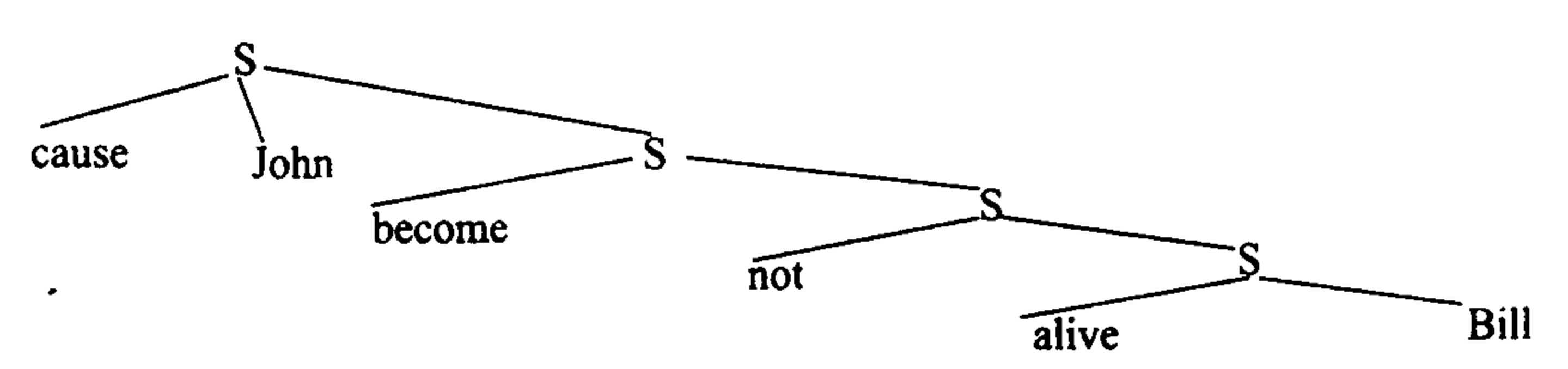


图 9.2.2.1 语义结构描写图

按照生成语义学的思路,以上 a、b、c、d 四个句子都可以用 d 来表示。我们可以用括号方式来描写 d 的构造:

$\langle_s \text{ cause John } \{_s \text{ become } (\text{not } [_s \text{ alive Bill}]) \} \rangle$

假设我们在词库中收集了与该问题相关的 3 个词项,并且得到了元语言的描写,以上各句的描写过程是:

词项	元语言描写	替换	所转换成的表层结构
dead	not alive	第 1 层 S	John caused Bill to became dead
die	become not alive	第 2 层 S	John caused Bill to die
kill	cause become not alive	第 3 层 S	John kill Bill
	树形图本身	不替换	John caused Bill to become not alive

不同的深层结构不过是选择了不同的词项。这种描写比原子命题的描写又深了一步。可以说,生成语义学在格语法的深度上把语义结构往深处推进了两次。第一次的推进由于是以原子命题为根据,并不分解词项,还比较容易达成共识标准,第二次的

推进由于开始用元语言分解词项,于是有两项工作必须得做:

1. 词典中的词项都必须经过元语言的描写。
2. 语义结构(或语义表达式,或谓词演算命题)必须按照元语言的描写来分析。

这两点在实际操作中会遇到一些困难,而且会带来很多随意性。

生成语义学处理语义比格语法更深。格语法有词汇的限制,动词和语义格都是词典中的单位,是千百年来在文化活动和经验活动中形成的结果。生成语义学没有词汇限制,而是用一些非常抽象的、独立于自然语言经验划分的元语言因子来说明语义关系。

格语法和生成语义学的深度可以通过“所有的人都要呼吸”看出来。在“格”语法中,“所有的人”是一个语义格,在生成语义学中,是一个谓词演算,即“人是所有的”,因此,生成语义学把格语法中的简单句描写成两个命题。这就出现了一个问题:生成语义学在取得共识标准上比格语法困难,比 Chomsky 的深层结构更困难。

9.3 深层结构中的词汇插入问题

Chomsky(1965)的深层结构假定相同语义的句子深层结构相同,不同语义的句子深层结构不同。从上面的分析看,生成语义学认为这两方面都有问题,这是取消深层结构的一个主要理由。Lakoff 和 Ross (1967)、Lakoff(1972)还从另一个角度提出了取消深层结构的理由,即有些词项插入必须在转换后进行。这里的论证过程在技术上相当复杂。概括地说,他们认为有些词项是否可以插入,需要首先知道是否可以转换,因此认为有些词项是在转换之后插入的。我们知道,Chomsky(1965)认为深

层结构是由改写规则和词汇插入确定下来的,转换是对这种已经存在的深层结构的操作。如果认为有些句子的生成需要先转换再插入词汇,就从一个角度否认了深层结构在转换之前就已经存在的结论,也就否认了深层结构的存在。Chomsky(1965, p103)实际上已经考虑到了这类问题,并且在严格次范畴化中把可转换和不可转换分成两类。Chomsky 所举的不可转换的例子是:

- John is resembled by Bill
- a good book is had by John
- John was married by Marry

Chomsky(1965, p104)指出:

That is, a Verb will appear in frame (56) and thus undergo the passive transformation only if it is positively specified, in the lexicon, for the strict subcategorization feature [— NP Manner,], in which case it will also take Manner Adverbials freely.

我们认为,有些词项是在转换前插入还转换后插入,是一个技术处理问题,并不是否定深层结构的理由。因此,尽管我们可以找出语义表达式存在的理由,但同时也可以找出深层结构存在的理由。这两个层面都有存在的必要性。再考虑汉语的一个实例。对于一个深层结构:

NP1 + V + NP2

现在来看转换和词汇插入的关系。如果在词库中选择“撕”,

就有:

NP1 + 撕 + NP2

如果要转换成“把”字句,首先需要知道 NP2 是有定的,而且动词前后必须有补充成分或限制形式。这时可以有两种办法,先转换成“把”字句:

NP1 + 把 + NP2 + V

再插入有定名词和限制动词,比如“撕了”:

张三 + 把 + 这本书 + 撕了

张三 + 把 + 李四写的这本书 + 撕了

按照这种方式,不仅词汇的插入可以在转换后,甚至短语结构规则的应用也可以在转换后进行:

NP→NP 的 NP (这本书→李四的这本书)

另一种办法是,先插入词汇,或进一步使用短语结构规则再插入词汇,由于最终目标是要生成“把”字句,所以只能插入有定名词和动词:

张三 + 撕了 + 这本书

张三 + 撕了 + 李四写的这本书

然后再转换成“把”字句。如果事先没有“把”字句的驱动,插入了不定名词:

张三 + 撕了 + 两本书

则不允许再转换成“把”字句。显然,以上两种生成过程都是可能的。无论选择哪一种生成过程,都必须要用到短语结构规则和转换规则。但在用转换规则前,已经有一个可以进行替换的结构,前面在讨论语义结构(论元结构)和深层结构的关系时,已经证明这个可替换结构不是词库中的格框架,而是深层结构。

生成语义学取消深层结构的另一个根据是转换会改变意义。按照 Chomsky(1965)的标准理论,动词和它所共现的名词构成的语义关系是句子的语义,被动、否定、疑问等是在这样的意义上再加上特定的意义,这就是深层结构需要解释的意义。但实际上量词、否定词在句子中的不同位置仍然会有不同的意义,这是深层结构所解释的意义没有包含的内容。从这一点说,生成语义学发现了另一个意义层面。而这个意义层面应该是语言学家可以解决的,并且应该解决的。但这并不构成否定深层结构的理由,因为我们同样可以在深层结构中标注量词、否定词所在位置的信息,就像在深层结构中标注被动、疑问等信息。

从根本上说,被动、疑问、否定、否定的位置、量词的位置等语义都不是语义格的语义,这些语义既可以在深层结构中标注,转换后不再改变意义,也可以不在深层结构中标注,通过转换来表示这些不同的意义,这是可以选择的,根本上是一个技术处理问题,20 世纪 70 年代转换生成语法在这个问题上的争论,并不是根本问题的争论。根本问题是转换和意义的关系,后面我们还要回到这个问题上来。

9.4 元语言问题

前面讨论的生成语义学的第 2 个表达深度本质上是按照分

析哲学中人工语言学派的方式研究语义问题。这需要建立很多元语言成分。这些元语言成分是怎么来的, 又是根据自然语言定义的, 因为自然语言是最初始的元语言。后期 Wittgenstein 就主张精确的描写自然语言。如果能解决好元语言问题, 仍然可能有效利用生成语义学的词项分解思想。下面结合元语言来分析问题。

9.4.1 自然语言是最初始的元语言

语义表达式是建立在谓词演算基础上的, 由于谓词演算有相对性, 语义表达式也有相对性。观察下面的例子:

有些北大学生学习美国英语

谓词描写 1:

设 $F(x, y)$: x 学 y

$G(x)$: x 是北大学生

$H(y)$: y 是美国英语

谓词描写 2:

设 $F(x, y)$, x 学 y

$A(x)$: x 是学生

$B(x)$: x 是北大的

$C(y)$: y 是英语

$E(y)$: y 是美国的

显然, 谓词描写 1 在结构上刻画得较粗, 但在词库中必须描写得更细, 词库中应该包括“北大学生”, “美国英语”这样一些片断。谓词描写 2 正好相反。生成语义学接近谓词描写 2, 格语法接近谓词描写 1。前面我们讨论的词项分解, 也可粗可细。这反

映了谓词描写的相对性。哪一种好是一个评价程序问题。不同的谓词描写方式各有自己的价值。

生成语义学必然要引起对元语言问题的追问。元语言问题是人工智能、语言学以及语言哲学的核心问题之一。从 Leibniz 到 Frege、Russell 以及维也纳学派,都希望设计一种人工的元语言来表述和分析哲学命题,这种元语言应该是高度精确的符号系统。这是一种理想观点。表述哲学的语言是否需要限制在人工语言中,取决于要解决什么问题。对于数学、逻辑学和很多自然科学来说,人工语言运动是很有意义的活动。但如果需要讨论的目标更广泛,包括语言学、哲学、美学、宗教等,人工语言就会面临很多困难。后期 Wittgenstein 以及 Moore、Ryle、Austin 则认为,日常语言(即自然语言)本身就是完善的,不必设计人工的元语言,哲学问题的表述和分析可以直接用日常语言,问题是哲学家首先必须描写清楚日常语言的意义。日常语言学派是把自然语言看成哲学表述的元语言,这是一个很有必要的转向。由于日常语言学派把自然语言看成是完善的,而词的意义主要是词的用法,所以精确描写自然语言中词的用法成为日常语言学派的首要工作。

所有的人工语言都是通过自然语言来定义的,人工语言肯定不能充分表述哲学命题,自然语言是最初始的元语言,于是元语言最终只能回到自然语言。我们说自然语言是最初始的元语言,这意味着不仅各层面的文化经验活动要由自然语言来组织,而且所有人工语言最终都要由自然语言来解释和定义。反过来,人工语言就很难解释和定义自然语言。有时候为了研究的方便我们需要一些形式化的语言符号来解释自然语言的一些重要概念,使自然语言的一些概念更明晰或加快思维的速度,比如符号逻辑 $\cup$ 、 $\cap$ 可以对自然语言的“或、且”作更严密的解释,微积分符号系统可以使运算更快、使思考更精密。日常语言中有一部分表达是不明确的,规定一种严格的人工语言可以使很多

表达明确化。这是人工语言的优势。比如:

自然语言: 85 分以上的算优秀

85 分本身算不算优秀的范围? 在很多自然语言中都有这种范围不明确的问题。

现在我们通过人工语言来表达

人工语言: 如果 $A > 85$, A 就是优秀

人工语言: 如果 $A > 85$ 且 $A = 85$, A 就是优秀。

这就是人工语言的好处。

但这些人工语言符号必须事先经过自然语言的解释和定义。既然自然语言能够定义出精确的人工语言, 这些精确的思想活动一定已经蕴含在自然语言中。比如上面的数学表达方式仍然可以用自然语言来表达:

自然语言表述: 85 分以上(不含 85 分)算优秀

自然语言表述: 85 分以上(含 85 分)算优秀

不过通常情况下, 人们对自然语言中蕴含了哪些复杂的观念并不清楚。人类各个层面的思维活动是否都依赖了自然语言, 现在还不太清楚, 但思想的表达和交流肯定要依赖自然语言中的范畴框架。

9.4.2 自然语言自身解释的内部循环

自然语言本身没有标记来区分对象语言的句子(对象语句)和元语言的句子(元语句), 因此自然语言中必然会存在哲学家经常提到的语义悖论。比如在黑板上写下这样一句话:

这句话是错的

无论断言这句话是真还是假,都会产生矛盾。如果断言这句话是真的,那就等于承认这句所断言的内容是真的,那这句话就是错的,因为这句话的内容是在断言这句话是错的。如果断言这句话是假的,就等于承认这句话没有错,那么这句话就是真的。人们对悖论的产生原因做出了各种不同的解释,从语言学的观点看,我们认为产生悖论的根本原因就在于没有区分对象语言和元语言,“这句话是错的”既是对象语言,又是元语言。

由于自然语言是最初始的元语言,解释和分析自然语言的规则和意义就只能靠自然语言本身,这就会造成自然语言自身解释的内部循环。自身解释内部循环是指:在一个自然语言系统中,如果充分系联各种定义和解释,最终会出现被定义项出现在定义项中。这种内部循环是违反形式逻辑和公理系统的,但在自然语言中是不可避免的。如果把自然语言中的某些语符看成是对象语言中的语符,另一些语符看成是元语言中的语符,表面上我们似乎可以用元语言中的语符来定义或解释对象语言中的语符:

对象语言	元语言因子 1	元语言因子 2		元语言因子 n
人	动物	会语言		
树	非动物	生物		
狗	动物	不会语言		
杀	引起	死		

从理论上说,我们似乎可以给每个对象语言中的语符一个完美

的元语言定义,一个语符不过是一束元语言因子的集合,但是这里的实质是容易看出来的,这只不过在对象语言和元语言之间建立了一套对应关系。拿表中对象语言的“人”来说,可以定义成“动物、有语言”,但“动物、有语言”这些元语言成分,并没有得到解释,如果再以汉语中的其他成分来定义“动物、有语言”,最终会进入内部循环。如果把自然语言中的规则分成对象语言规则和元语言规则,那么元语言规则通过什么解释?也会出现规则解释的内部循环。

自然语言自身解释的内部循环性可以在词典释义中观察到,比如商务印书馆 1996 年版《现代汉语词典》所作的解释:

高:从下向上距离大

上:位置在高处的

下:位置在低处的

不:用在动词、形容词和其他副词前面表示否定

否:表示不同意

是:对,正确

对:相合,正确,正常

正确:符合事实

符合:相合

以上定义均存在间接循环。通过系联很容易观察到间接循环本质上也是直接循环。比如上面第一对实例,“高”的定义项中有“上、下”这两个词,把这两个词的定义代入“高”的定义,就有:

高:从下向上距离大→高:从位置在低处向位置在高处距离大

通过这种代入,可以发现定义中包含了“高”这样的被定义项。

被定义项和定义之间的根本关系就是对象语言和元语言之间的关系,既然词典中可以找到大量被定义项被包含在定义项中的实例,说明词典不能有效的区分开对象语言和元语言,或者说词典中包含了自然语言解释的内部循环。语法手册的道理在这方面和词典一样,也包含了自然语言规则在解释上的内部循环。

回到根本问题上来。生成语义学派想把自然语言语句的解释基础建立在语义表达式上,并进一步通过词项分解(lexical decomposition)来描述语义表达式。McCawley (1968)认为下面语句的语义表达式是相同的:

- a. John killed Bill
- b. John caused Bill to die
- c. John caused Bill to become dead
- d. John caused Bill to become not alive

McCawley 的做法实际上也涉及了元语言问题,因为哪些词项不可以分解,哪些词项可以分解,本质上是在规定对象语言和元语言。我们可以把以上各句的元语言解释过程表述如下:

对象语言	元语言描写方式
kill	cause to die
die	become dead
dead	not alive

这里的对象语言和元语言关系是有层次的,kill只是对象语言中的语符,所以只在对象语言一栏中出现,die 和 dead 既是对象

语言中的语符,也是元语言中的语符,所以在对象语言和元语言两栏中都出现,alive、cause、become、to 则都只是元语言,只在元语言一栏中出现。类似 alive 这样的元语言语符,怎样解释和定义?如果我们再进一步查阅《牛津高级英汉双解词典》,就会发现 alive 的解释是 not dead,这里又出现了内部循环。可见,生成语义学的语义表达式建立在哪个语义层面上,用哪些语符作语义表达式,哪些语符可以分解,哪些不可以分解,都不能回避内部循环问题。

波兰逻辑学家 A. Tarski(塔斯基,1949)认为语义悖论可以通过抽象语言来解决。他认为语言是分很多级的。我们再用上面的“这句话是错的”这一实例来分析塔斯基的观点。如果“这句话”指的是“雪是白的”,我们可以有:

A:雪是白的

B:A 是错的

C:B 是错的

A 是对象语句,B 是 A 的元语句,C 又是 B 的元语句。通过这种语义层级的区分,就不会产生悖论。Tarski 的语言层级思想能够解释语义悖论产生的根源,但不能否认自然语言自身解释的内部循环问题,因为要对 C 语句有所断言又会追问到更多的抽象语句,最终形成自然语言解释的内部循环。

9.4.3 获得元语言的可能性:核心语符和核心规则

尽管对自然语言的解释存在内部循环,但有文字的民族学习语言确实充分利用了词典和语法手册,而对词典和语法手册的描写借助的都是自然语言,因此自然语言正在被自然语言陈述着,自然语言释义的内部循环性仍然是有价值的。

那么人们解释或定义自然语言中的一个语符或一条语法规

则是靠什么呢？哲学家解释或定义一个概念或一条陈述是靠什么呢？应该是靠母语中的核心语符和核心规则。人类把自然语言作为母语来学习，首先需要掌握核心语符和核心规则，然后通过核心语符和核心规则定义和解释非核心语符和非核心规则。获得核心语符和核心规则这个过程不是通过定义和解释来完成的，即人类不是通过某种比自然语言更初始的元语言来学习自然语言。人类掌握母语中的核心语符和核心规则有两个基本过程：

1. 获得一批言语片断的指称和用法。
2. 从这些言语片断中提取有限的语符和规则。

第一个过程是一种反复刺激的过程，儿童不断听到“老鼠、老鹰”这样一些和生活行为最密切的言语片断，不断看到这些片断所指称的事物和事件，就获得了这些言语片断的语音形式以及这些语音形式所指称的意义。这个过程和 Bloomfield(1933)提到的行为刺激理论以及 Wittgenstein(1953)所提到的意义即用法有关系。Wittgenstein 说：教孩子说话靠的不是解释或定义，而是训练。但 Bloomfield 和 Wittgenstein 重点在解释意义是什么。从元语言来源的角度看，我们认为意义来源不仅和用法有关系，而且和指称有关系(陈保亚, 2006, 1)，即这些言语片断的意义是在反复使用和指称过程中形成的，而不是通过定义或解释形成的。

通过分布和指称获得部分言语片断的意义，这只是言语习得的第一步。母语者不可能记住每一个言语片断的用法，因此仅仅获得部分言语片断的用法是不够的。自然语言不是若干言语片断的堆积，而是通过有限的规则和语符生成无限的言语片断(Chomsky, 1965)。当儿童通过第一个过程获得若干核心言语片断后，就需要通过一些方式获取有限的语符和规则，这就是

单位和规则的还原过程。

我们根据这些年来对儿童学习语言的一些初步记录,尤其是对儿童陈赅的语言学习的长期记录注意到儿童获得母语中核心语符和核心规则的最基本方式并不是通过定义或解释,而是通过观察部分言语片断的用法(分布)和指称以及对言语片断进行平行周遍对比。人类习得母语核心语符和核心规则的过程或许还有一些其他的方法,或依赖了其他一些条件,但肯定不是通过其他元语言的定义和解释,这是一个可以观察到的事实。

在自然语言内部循环解释的过程中,人们可以通过核心语符和核心规则定义非核心语符和非核心规则,核心语符和核心规则构成了自然语言中的元语言核心,而核心语符和核心规则的获得最终是通过言语片断的用法(分布)和指称以及对这些片断的平行周遍对比。至于哪些是核心语符和核心规则,哪些是非核心语符和非核心规则,一般情况下容易体会到,但如果要追问一个严格的条件或界限,现在还没有办法回答,这是一个需要深入研究的问题,可能涉及语言认识论中的某些根本原则。

核心单位和核心规则的获取来自于平行周遍对比,而平行周遍对比从根本上说建立在相似性基础上,这就会伴随不完备性。不过自然语言也不像人工语言学派认为的那样,可以被人工语言取代。自然语言是最初始的元语言,这是自然语言无所不表述的强大功能决定的,也是自然语言核心规则和核心语符的获取方式决定的。Wittgenstein(1953)把自然语言比成游戏, Saussure(1916)把自然语言比成符号系统,从元语言的角度看,这些比较都不是很贴切的,都忽视了自然语言的初始性,因为自然语言以外的任何一种游戏或符号系统被禁止后,人们都可以继续思考和交流,而一旦自然语言被终止或抛弃,人们将无法思考和交流。语言比游戏和符号系统更为初始。语言哲学家应该充分比较不同的自然语言,深入分析自然语言,确定核心规则和核心语符,充分描写核心规则和核心语符,比较不同语言之间核

心规则和核心语符的差异,找出完备性和不完备性,在此基础上构建一种表述功能完备的、基于自然语言的元语言。从这个角度看,人工语言学派和日常语言学派都有合理的价值。语言分析不应该仅仅限于日常语言学派所做的语符分析,还应该包括规则的分析,语言哲学的语言分析根本目标不是精确性分析,而是完备性分析。通过完备性分析充分认识自然语言在哪些方面是完备的,哪些方面是不完备的,哲学家和语言学家为了建构一种完备的元语言需要在哪些方面对自然语言进行人工补充。在分析和建构方面,对自然语言的分析更优先,因为自然语言是最初始的,从自然语言中获得的元语言规则更容易被母语者接受,更反映母语者认识活动的原型。

9.4.4 核心规则的建构分析

建立元语言需要做大量核心规则的分析,这不是不可能的。下面尝试做一个核心规则的分析。

在逻辑学、哲学乃至一般科学活动中,概念分类和整体切分是两种非常基本的认识活动,概念分类是把一个属概念划分成两个或几个种概念,或者说是把大类划分成小类。但表达这两种认识活动的词语经常互用。下面的陈述无论正确与否,都是在进行概念分类:

语言哲学可以划分成语言分析哲学和语言解释哲学
语言学可以划分成理论语言学和应用语言学
哲学可以划分成东方哲学和西方哲学
词可以划分成实词和虚词

整体的切分是把整体切分为部分。下面的陈述都是在进行切分:

命题可以划分成主词和宾词

地球可以划分成南半球和北半球

句子可以划分成主语和谓语

汉语的音节可以划分成声母和韵母

以上两类陈述都用了“划分”这个词。从 Wittgenstein 或日常语言学派的角度看,有人可以对“划分”一词的用法作充分的分析,直到发现这个词至少有“概念的分类”和“整体的切分”两种用法。但是在具体的命题中,不是所有的人都很容易判定下面的实例是在分类还是切分:

人可以划分成男人和女人

句子可以划分成陈述句、疑问句、感叹句和祈使句

句子可以划分成主语和谓语

四边形可以划分成两个三角形

四边形可以划分成平行四边形和不平行四边形(指分类)

这需要根据上下文来判定是哪一种情况。需要靠上下文来领悟。领悟不是判定歧义的有效方法。由此可见,弄清一个词的用法是一回事,判定一个命题的含义是另一回事。后期 Wittgenstein 和日常语言学派重点讨论了前一项工作,很少涉及后一项工作。要找出自然语言中的核心规则,必须要讨论后一项工作。

有人可能会想到人为地给“划分”和“切分”指派特定的含义,“划分”只表示从大类到小类的分类,“切分”只表示从整体到局部的分割。但由于自然语言是约定俗成的系统,支持和维护这种约定俗成规则的是使用自然语言的庞大群体,给自然语言中的词指派特定语义的活动很难实现。在自然语言基础上建构

元语言的一个首要原则是：

不要给自然语言的词或规则任意指派超出约定俗成范围的语义。

如果没有一个形式判定标准，从属到种的概念分类与从整体到局部的切分很容易混淆，而人们对“划分”、“分成”这几个词语的混用又加深了这种混淆。

在认识活动中，外延的扩大或缩小，或概念的限制与扩大，也是一项基本认识活动，这一活动和概念的分类是同一个问题的两个角度，在概念分类缺少形式判定标准的情况下，判定以下关系是外延扩大还是范围扩大也会出现问题：

中国→亚洲

诗歌→文学形式

语法学→语言学

珠穆朗玛峰→喜马拉雅山

认识活动中最常用的“属于”一词，不能把外延扩大和范围扩大两种情况区别开：

中国属于亚洲

诗歌属于文学形式

语法学属于语言学

珠穆朗玛峰属于喜马拉雅山

无论是概念分类也好，外延的限制或扩大也好，本质上都需要明确大类和小类的关系。现在的问题是，汉语中有没有判定概念分类的表达式？有没有判定外延的限制或扩大的表达式？

先考虑一种替换分析(带星号*的句子是替换不成功的):

人可以划分成男人和女人

→* 人可以切分成男人和女人

句子可以划分成陈述句、疑问句、感叹句和祈使句

→* 句子可以切分成陈述句、疑问句、感叹句和祈使句

句子可以划分成主语和谓语

→ 句子可以切分成主语和谓语

四边形可以划分成两个三角形

→ 四边形可以切分成两个三角形

四边形可以划分成平行四边形和不平行四边形

→* 四边形可以切分成平行四边形和不平行四边形

凡是能用“切分”替换“划分”的句子是表示从整体到局部的切分的句子。但是,不能用“切分”替换的,不一定不是表示整体到局部的切分。比如:

欧洲可以划分成东欧和西欧

→* 欧洲可以切分成东欧和西欧

仔细领悟,原句本来是表示整体到局部的切分,但不能用“切分”替换。可见,虽然能用“切分”替换的句子的表示从整体到局部的切分,但不能用“切分”替换的句子不一定表示分类。一般地说,用语词替换来锁定一个句子的意义并不具有充分性。

再考虑通过转换分析来建构范畴表达式。Chomsky(1957)首先提出转换是生成句子的必要规则,并深入研究了这种转换,后来又用移位(Chomsky, 1981)来表述转换。我们不打算讨论

转换和生成句子的问题,下面我们重点考虑怎样通过转换建构范畴表达式。自然语言的判断句恰好有一类是表示小类和大类关系的:

日常语言学派是分析哲学

这里主语“日常语言学派”是小类,判断词“是”后面的宾语是大类。一般地说,概念分类都可以转换成另一种判断句:

分析哲学可以划分成逻辑经验主义和日常语言学派

→逻辑经验主义是分析哲学

→日常语言学派是分析哲学

概括成一般格式是:

NP 可以划分成 NP1 和 NP2

→NP1 是 NP

→NP2 是 NP

于是我们有可能将自然语言中的判断句作为确定大类和小类关系的表达式,不过这种可能性还不能马上实现,因为以上转换格式是不可逆的,所有的概念分类句都可以转换成判断句,但不是所有的判断句都可以逆转换成概念分类句:

儿童是花朵,儿童是希望

→*儿童可以划分成花朵和希望

更严格地说,概念分类句转换成判断句可以周遍到每个实例,但判断句逆转换成概念分类句不会周遍到每个实例。转换不可逆

的原因在于,汉语的判断句不仅可以表示大类和小类之间的关系,还可以表示同一、个体集合、明喻以及特定语境下的自由关联等关系:

《工具论》的作者是亚里士多德(主语和宾语是同一关系)

亚里士多德是哲学家(主语和宾语是个体和集合的关系)

眼睛是心灵的窗户(主语和宾语是隐喻关系)

他是日本太太(特定语境自由关联:他的太太是日本人)

判断句的种属判定功能、同一性判定功能、个体集合判定功能、明喻判定功能以及特定语境自由关联功能都没有区别性形式标记,而我们需要的是种属判定功能的判断句。考虑到转换的不可逆关系,引入判断句作概念分类的表达式需要加以限制,以便和其他表达功能的判断句分开。我们在判断词“是”和宾语之间加上一个形式标记“一种”,这样就构成了既具有形式区别特征而又表示种属关系的判断句:

S: NP1 是一种 NP2

这一格式能够把同一性判断、个体集合判断、隐喻判断以及自由关联判断排除,比如:

同一关系(前面加星号*的句子是不可接受的句子):

《工具论》的作者是亚里士多德→*《工具论》的作者是一种亚里士多德

亚里士多德是《工具论》的作者→*亚里士多德是一种

《工具论》的作者

相对论的创始人是爱因斯坦→* 相对论的创始人是一种爱因斯坦

爱因斯坦是相对论的创始人→* 爱因斯坦是一种相对论的创始人

个体集合关系(以下换成“一位”成立,但不是种属关系):

亚里士多德是哲学家→* 亚里士多德是一种哲学家

高斯是数学家→* 高斯是一种数学家

海明威是文学家→* 海明威是一种文学家

爱因斯坦是物理学家→* 爱因斯坦是一种物理学家

隐喻关系:

眼睛是心灵的窗户→* 眼睛是一种心灵的窗户

儿童是花朵→* 儿童是一种花朵

指挥是乐队的灵魂→* 指挥是一种乐队的灵魂

数学是科学的皇后→* 数学是一种科学的皇后

特定语境自由关联关系:

他是日本太太→* 他是一种日本太太

他是第二名→* 他是一种第二名

他是大眼睛→* 他是一种大眼睛

他是皮鞋→* 他是一种皮鞋

在“NP1 是一种 NP2”判断句中, NP1 是小类(种), NP2 是大类(属)。这正是我们需要的判定式。如果是概念分类, 子项和母项就存在种属关系, 即小类和大类的关系, 子项就可以进入 NP1, 母项就可以进入 NP2; 否则就不是概念分类, 而是从整体到局部的切分。再来看我们前面提到的例子(见表 9.4.4.1):

表 9.4.4.1 分类和切分的表达式判定

分类或切分陈述	[NP1 是 一 种 NP1]的判定	认识活动的 性质
句子可以划分成主 语和谓语	· 主语(或谓语)是一 种句子	切分
地球可以划分成南 半球和北半球	· 南半球(或北半球) 是一种地球	切分
命题可以划分成主 词和宾词	· 主词(或宾词)是一 种命题	切分
人可以划分成男人 和女人	男人(或女人)是一 种人	种属划分
句子可以划分成陈 述句和非陈述句	陈述句(或非陈述 句)是一种句子	种属划分
词可以划分成实词 和虚词	实词(或虚词)是一 种词	种属划分

Wittgenstein 在《哲学研究》(1953)中提到过两个判断句：

玫瑰是红的
二加二是四

Wittgenstein 认为这两个“是”不同，因为第一个句子中的“是”不能用“等于”去替换。从更高的层次看，根据我们给出的“NP1 是一种 NP2”，可以判定这两个句子都不涉及概念分类。再来看外延的扩大和限制问题(见表 9.4.4.2)：

表 9.4.4.2 外延关系的表达式判定

是否有扩大和限制关系	[NP1 是一种 NP1]的判定	认识活动的性质
中国- 亚洲	· 中国是一种亚洲	局部与整体
袖子- 衣服	· 袖子是一种衣服	局部与整体
珠穆朗玛峰- 喜马拉雅山	· 珠穆朗玛峰是一种喜马拉雅山	局部与整体
诗歌- 文学形式	诗歌是一种文学形式	限制与扩大
微积分- 数学	微积分是一种数学	限制与扩大
植物- 生物	植物是一种生物	限制与扩大

这时判定式“NP1 是一种 NP2”仍然有效。“诗歌”和“文学形式”分别进入 NP1 和 NP2 以后,判定式“NP1 是一种 NP2”在自然语言意义上是可接受的,则“诗歌”到“文学形式”是外延扩大。“中国”和“亚洲”分别进入“NP1 和 NP2”以后,判断式“NP1 是一种 NP2”在自然语言意义上不可接受,不是外延扩大,而是范围扩大。与此不同,上面的 6 个实例中的 NP1 和 NP2 都可以进入表达式“[NP1 属于 NP2]”,由此也可以证明“[NP1 属于 NP2]”表达了外延扩大和范围扩大两种完全不同的认识范畴,可以作为一种核心语义表达式。

9.5 语义表达式的初始性

Chomsky 的深层结构决定了语义解释,因此 Chosmky 的语义解释实际上包括深层结构中词的意义和词的语法关系。Fillmore 在语义格的关系上建立深层结构,因此他的语义解释包括句子中的词和词的语义格关系。比较两者的异同,深层结构和格语法的语义解释的根本区别在于语法关系和语义格关系。如果按照 Chomsky 的说法,格关系已经包含在语义次范畴

中,则格语法在解释语义上和深层结构没有根本的区别。

生成语义学由于对词项进行分解,所以在解释语义上和标准理论、格语法都有比较大的区别。以“曹操杀了关羽”为例,三种语义解释是:

深层结构语义表达式: $\langle_s(\text{NP 曹操})(\text{VP}[\text{V 杀了}][\text{NP 关羽}])\rangle$

格语法语义表达式: $\langle_s(\text{NP 曹操}^A)(\text{VP}[\text{V 杀了}][\text{NP 关羽}^0])\rangle$

生成语义学语义表达式: $\langle_s(\text{NP 曹操})(\text{VP}\{\text{V 让}\}\{\langle_s[\text{NP 关羽}][\text{VP 死了}]\rangle\})\rangle$

以上右上方的标注表示语义格。前面我们已经分析过,生成语义学的词汇分解语义表达式存在两个问题:

1. 被分解项和分解式的语义对等问题
2. 元语言问题

不过,词典中的“杀”仍然需要语义解释,从这个角度看,生成语义学实际上是把词典中要做的工作的一部分放到语义表达式中了。由于语义表达式是语法的一个部分,这就等于把词典解释的一个部分放在句法中了。

在语符库中仅仅给出语符 V 的语义格是不够的。比较:

张三买了我们一亩地

张三卖了我们一亩地

语符库中语义格的标注是相同的:

买 [施事、与事、受事]

卖 [施事、与事、受事]

而且这里的“与事”和语符“给”的“与事”并不完全相同。在语义格之外,还有受事移动方向的问题:

买 [施事、与事、受事], “受事”从“与事”转移到“施事”

卖 [施事、与事、受事], “受事”从“施事”转移到“与事”

如果有了移动方向,就可以更进一步解释一些现象。像下面的语符 V:

张三借了我们一亩地

张三租了我们一亩地

这里的语义结构解释和“买、卖”又不同:

借 [施事、与事、受事], “受事”从“与事”转移到“施事”,或者从“施事”转移到“与事”

租 [施事、与事、受事], “受事”从“施事”转移到“与事”,或者从“与事”转移到“施事”

这样的方向问题还存在于其他语义结构的动词中,比如:

今天下午上课

→今天下午听课

→今天下午讲课

以上问题都是语义表达式需要解决的问题。语义表达式还可以

解决深层结构和语义格不能解决的另一个问题。在语义格中, 语义组合的单位是词或以词为中心的词组。比如:

人能够思维

但是有些句子“施事”或“格”不是直接取自词典中的词, 而是通过陈述转换过来:

会语言的能够思维

“会语言的”是词库中找不到的临时组合, 我们能够将其和“人”对等起来, 或者说用“会语言的”去替换“人”, 本身说明了我们在进行比语义格更深的语义分析。

从语言生成的角度看, 在经验活动过程中, 首先是一些语义观念的形成, 然后才有表达的需要。语义表达式如果尽可能的接近语义观念, 就能说明句子的生成过程。

从语言翻译的角度看, 两种语言不仅没有完全相同的表层结构, 甚至表达深层结构的改写规则(按照 Chomsky 所定义的)也不同, 语义格也不同, 词库也不能完全对接, 因此, 翻译过程如果建立在深层结构基础上, 会遇到困难。如果能找到普遍的语义表达式, 就可能对代码转换、翻译作出有价值的贡献。比如:

John lent Jack a book

对译成汉语是:

约翰借给杰克一本书

汉语词库中通常是不会收“借给”这个组合的, 这样翻译 lend 实

际上是把 lend 的语义分解成了“借”和“给”两个义项。

当然,词典的语义解释并不是规则信息,而是个别事实,本质上仍然需要记忆,所以在句法中描述词项的语义分解式并不能简化描写。不过词典中如果能够给出对等词项的语义分解式,对于说明深层结构不同语义解释相同和深层结构相同语义解释不同是有价值的。

回到语义解释的层面上来。深层结构、语义格、语义表达式都是必要的。建立深层结构不仅可以使描写简单化,即有些句式用短语结构描写,有些句式在转换规则上处理,同时可以看出句式之间的关系;更重要的是,很多线性结构是不可观察的,如“他[正在][在书房]读书”。格语法能够解释深层结构不能解决的语义功能问题,即有些歧义或语法关系必须追问语义格的性质才能说清楚。生成语义学的语义表达式尽管还没有达成共识,但任何一种语符库中都需要用一些更核心的语符和规则来表述另一些次核心的语符和规则,这本质上就是语义表达式的工作。

思考问题

1. 如何解决元语言的内部循环问题?
2. 如何利用词汇分解进行句法研究?

10. 语义解释层面

10.1 语义格框架和投射原则

20 世纪 60 年代的标准理论、格语法、生成语义学,分别在深层结构、格框架、语义表达式上解释语义,从根本上看是想在深层结构解释语义,只是深层结构的深度不一样。根据这种思想,转换只是把句子的深层结构转换成表层结构,表层结构通过语音规则得到可感知的语音序列。这就是著名的“转换不改变意义”的思想。

Chomsky(1965,p136)说明深层结构解释语义的实质,语义解释仅仅依赖词汇项目、语法功能和语法关系:

Thus when we define “deep structures” as “structures generated by the base component,” we are, in effect, assuming that the semantic interpretation of a sentence depends only on its lexical items and the grammatical functions and relations represented in the underlying structures in which they appear. This is the basic idea that has motivated the theory of transformational grammar since its inception.

Chomsky(1965,p136)还解释了上述在底层结构(underlying structure)解释语义的观点来源于 Katz and Fodor(1963)和 Katz and Postal(1964):

Its first relatively clear formulation is in Katz and Fodor(1963), and an improved version is given in Katz and Postal(1964), in terms of the modification of syntactic theory proposed there and briefly discussed earlier.

不过语法功能、语法关系在这里的含义不是很清楚的。考虑下面实例：

狗咬猫→猫被狗咬了

狗咬猫→咬猫的狗

狗咬猫→猫咬狗

第1句原来的宾语“猫”变成了被动句的主语，表层的语法功能和语法关系显然变了，但通常被看成是没有改变意义的转换。这个时候“意义”指什么？第2句由一个句子变成了一个名词性结构，表层的语法功能和语法关系也改变了。如果说这个时候意义没有改变，“意义”又指什么。第3句主语和宾语交换了位置，语法功能和语法关系显然变了，这时的意义一般都同意已经改变，这时的意义又指什么？根据 Harris(1952)和 Chomsky(1957)的用例看，第1、2个实例的变化可以看成转换，第2个变化通常不会被看成转换。为什么？

其实语言结构可以有各种不同的变形方式：

警察打了司机

1. 司机被警察打了
2. 警察把司机打了
3. 打司机的警察
4. 警察打的司机
5. 警察打的是司机

6. 打司机的是警察
7. 是警察打的司机
8. 是打司机的警察
9. 司机打了警察

第 1、2 句很容易被看成是原句“警察打了司机”的转换。第 9 句是否是原句的转换？由于施事和受事的地位已经发生了改变，就不能看成是原句的转换。第 3 到 8 句是否是原句的转化？尽管意义差别很大，3、4 是名词性的，其他句子则是“是”字结构，表示强调，但都保持了施事和受事关系不变。我们说，只要语义结构关系保持不变的一系列变式，都是转换的结构。这样的规定对语义结构的理解提供了一个方便。其实这样的规定正是 Chomsky (1965) 投射原则 (projection principle) 的实质。Chomsky (1981, p29) 进一步给出了投射原则：

Representations at each syntactic level (i. e. , LF, and D- and S-structure) are projected from the lexicon, in that that they observe the subcategorization properties of lexical items.

考虑我们曾经分析过的 Chomsky 规则部分的一个实例(汉语版 p42—43, 英文版 p33), 来说明投射原则的情况：

(48) it is unclear who to see

(49) (i) it is unclear [_S who [_S to see]] (表层结构, 和可观察的句子基本是一致的。)

(ii) it is unclear [_S whoi [_S PRO to see t_i]]
(S-结构, 有语迹)

(iii) it is unclear [_S COMP [_S PRO to see who]]

(D-结构,尚未经过转换,所以不存在语迹)

(iv) it is unclear [_S for which person x [_S PRO to see x]] (LF,即逻辑形式,是解释性的,由 S-结构派生出来的)

该实例比较详细地说明了各表达层次的关系。由于 iv 是从 ii 派生出来的,因而逻辑形式是解释性。其中 LF、D-structure、S-structure 都要遵守次范畴特征,即要从词库得到投射,或者说词项的次范畴特征要投射到这 3 个表达式上。trace i (t_i) 的确定就是由 projection principle 决定的。在词库中,see 要有 direct object。see 词项的次范畴特征是:

NP-see-NP

S-结构(ii)中,see 前面的 PRO 表示 see 实际上有一个施事,see 后面的 t_i 表示 see 实际上有一个受事,现在都标注出来,以此满足投射原则。D-结构(iii)中,see 后面用 who 表示底层宾语,即这个成分还没有移位,因此也无所谓留下 trace。逻辑式(iv)是在 S-结构上派生出语义解释,也保留了词项中的次范畴特征。在以上关系中,Chomsky 的次范畴特征也相当于语义结构。又比如(Chomsky,1981,p54):

(14) John was [_a killed]

(15) John INFL be [_a kill' t] (S-structure)

(16) [_{NP} e] INFL be [_a kill John] (D-structure)

(15)是(14)的 S structure,这里的 kill' 是 kill 的一种形式,在这里是被动表达一种方式,由于受事 John 移动到前面,留下语迹 t。(16)是 D structure,深层宾语直接用 John 标注。无

论在 S-结构还是 D-结构中,都保持了相同的语义结构关系。以上的 INFL 是指屈折成分(inflection),是句子 S 的核心,句子被看成是屈折短语(IP),这是支配与约束理论看待和分析句子的一种观点。

再比如汉语实例:

汪锋把杯子砸了

DS: 汪锋把_[NP e]砸了杯子

SS: 汪锋把_{[杯子]_i}砸了 _{t_i}

LP: 汪锋把 X 砸了 X

相当于表层结构的 SS 要解释语义,就得有语迹标注。

Chomsky(1981)对投射原则的解释是以词项的次范畴化特征(the subcategorization properties of lexical items)为基础的。我们也可以从论元的角度来理解投射原则,即在各种转换式中,论元角色不变。比如:

汪锋^{施事}砸了杯子^{受事}

→汪锋^{施事}把杯子^{受事}砸了

→杯子^{受事}被汪锋^{施事}砸了

→杯子^{受事}汪锋^{施事}砸了

→砸杯子^{受事}的是汪锋^{施事}

→是汪锋^{施事}砸的杯子^{受事}

→汪锋^{施事}砸的是杯子^{受事}

→汪锋^{施事}砸的杯子^{受事}

→砸杯子^{受事}的汪锋^{施事}

无论怎样转换,“汪锋”总是施事,杯子总是受事。以上“汪锋砸的杯子”有句子和名词短语两种歧义,最后一个例子“砸杯子的

汪锋”是名词短语,但都不改变“汪锋”作为施事、“杯子”作为受事的论元关系。这就是投射原则的本质。

前面说,按照投射原则,各个表达层面的语义结构都要保持语义结构中的必要论元(argument)或配价。不过在各种转换模式中,陈述式和指称式有较大的区别,陈述式对必要论元的要求很高,缺必要论元会不符合语法,指称式对必要论元的要求略低一些:

孩子理解音乐
理解音乐
* 孩子理解

孩子对音乐的理解
孩子的理解
对音乐的理解

在英语中,陈述式对论元的要求更严,逻辑主语通常也不能隐藏:

The students translated the book
* translated the book
* The students translated

the students' translation of the book
the students' translation
the translation of the book

10.2 语义结构的语义和转换的语义

在保持语义结构关系不变的前提下,或者说在满足投射原

则的前提下,转换是否会改变语义?

自从 Chomsky 的 Syntactic Structure(1957)发表以来,很多学者对转换做了研究,其中的一个核心问题是转换是否要处理语义问题。先让我们再来进一步分析转换完全不改变语义的含义。Chomsky(1957)的转换就有一部分不涉及意义的改变:

John past eat an apple→John ate an apple

这部分转换 Chomsky 称为必选转换(obligatory stransformation)。这部分句子的生成特点是把抽象的结构实现为可观察结构。为了更好地理解下面的问题,我们认为这里的转换的另一个本质特点是:从抽象结构到可观察结构,只有一种转换。

但是被动句的转换和以上转换不一样。被动句的转换通常被称为任选转换(optional transformation),和主动句相比,这类转换已经改变了表层的语法功能和语法关系。Hockett(1961)认为被动转换应该在底层标注,这已经有了转换不涉及意义的思想。有的学者主张转换不改变语义,认为 Chomsky(1955, 1957, 1962)中的任选单一转换(optional singulary transformation),比如被动转换、疑问句转换等,必须表述成必选转换(obligatory transformation)。确定这些转换的使用是在语符串中加某个标记,比如被动句标注 by passive。这方面有代表的研究有 Lees(1960)对否定转换(negation transformation)的研究、Fillmore(1963)对嵌入转换(embedding transformation)的研究、Fraser(1963)对连接转换(conjoining transformation)的研究、Klima(1964)对疑问句的研究。Katz 和 fodor(1963)首先表述了一种思想,句子的语义解释取决于词汇项、语法功能和语法关系,这些内容都在深层结构中,Katz 和 Postal(1964)进一步把这一思想系统化了,并提出了一个很重要的原则:

the only contribution of transformations to semantic interpretation is that they interrelate Phrase-markers.

这一原则又叫 Katz-Postal principle。这一原则的含义是,转换除了连接短语标记,对语义解释并不起作用。因此转换不能引入带有意义的成分,也不能删除不能恢复的词汇项目。

Katz 和 Postal(1964)提出的理论的核心思想是转换不改变意义。但实际上前面几个可选转换的实例中,箭头左右都有不同性质的意义的改变,只不过不是论元关系的改变。因此,转换是否改变意义取决于我们怎样理解意义。从转换规则研究的文献看,转换规则涉及的主要是否定转换、疑问转换、被动转换、名物化转化、嵌入转换等,这些转换都没有影响语义结构关系。比如:

NP1 + V + NP2

→ The cat caught the mouse

→ The cat didn't catch the mouse (否定转换)

→ Did the cat catch the mouse? (一般疑问句转换)

→ The mouse was caught by the cat (被动转换)

→ the catching of the mouse (名物化转换)

→ the catching of the cat (名物化转换)

其中名物化是有歧义的,但只有和原来的语义结构关系不冲突的名物化转换才是合理的转换,因此, Katz 和 Postal 转换不改变意义的理论可以概括为转换不改变语义结构关系。Chomsky (1965)所说的转换不改变语法结构关系和语法功能,应该理解成深层结构中的语法关系和语法功能。上面“猫追老鼠”变形为“老鼠追猫”,改变了语义结构关系,即改变了题元关系,已经不属于转换。实际上,从 Harris(1952)和 Chomsky(1957)开始,

转换的实例都没有改变语义结构关系,当时这一点并没有得到明确表述,主要是因为当时的描写都不打算涉及语义。

Chomsky(1965)开始涉及语义,吸收了 Katz 和 Postal 的转换不改变意义的主张,把任选转换部分因为转换引起的句式语义改变的信息都放在基础部分,比较主动句和被动句的转换:

$$NP1 + V + NP2 \rightarrow NP1 - V - NP2$$

$$NP1 + V + NP2(\text{by passive}) \rightarrow NP2 - \text{be} - V - \text{ed} - \text{by} - NP2$$

即在深层结构上给出被动标记,或者说把“被动”的意义移到深层结构上,转换过程中不再改变意义。实际上这只是一种技术处理方案,把表层句式的差别解释为深层结构上的差别。但是,即使我们把转换引起变化的信息放到深层结构中处理或者说标注,我们仍然面临这样的问题:“狗咬猫”和“猫被狗咬了”属于不同的句式,通常都认为有“主动”和“被动”的语义差别,这种差别语义有没有改变?

有一点是清楚的,主动和被动的差异是平行的,比如:

John broke the cup $\rightarrow$ The cup was broken by John

John ate the apple $\rightarrow$ The apple was eaten by John

又比如汉语的例子:

张三砸了杯子 $\rightarrow$ 杯子被张三砸了

张三喝了酒 $\rightarrow$ 酒被张三喝了

张三卖了房子 $\rightarrow$ 房子被张三卖了

换个角度说,左边的句子在“主动”的语义上是平行的,右边句子在“被动”的语义上是平行的,箭头两端的语义结构关系是相同

的,朱德熙(1986)讨论过这种平行现象。我们的问题是,如果把这种平行关系看成是意义相同,那么转换不改变意义就是指不改变语义结构关系或论元结构关系。如果认为“主动”和“被动”意义不同,那么转换会改变意义。这时的意义相同不仅要求题元关系相同,而且要求高层次“主动、被动”等意义也相同,即要求句式义相同。

从以上分析可以看出,转换是否改变意义在某些方面是没有争议的,比如必选转换,在另一些方面是有争议的,比如任选转换,这取决于我们怎么理解句子意义不变的含义。

我们认为转换不改变意义的提法是有问题的,因此在对语言结构的认识上产生了混乱。严格地说,各种转换形成的不同表层结构都有意义上的差别,它们只是在语义结构关系上相同:

汪锋砸了杯子
杯子被汪锋砸了
汪锋把杯子砸了
杯子汪锋砸了

这些看上去可以解释成意义不变的句式,在语篇连接上是有差别的:

谁砸了杯子:
→汪锋砸了杯子
→*杯子汪锋砸了。
→*杯子被汪锋砸了

这些语义差别都是通过表层结构体现出来的。如果把句子的意义仅仅理解成语义结构关系中包含的意义和具体的词的词典意义,则转换确实没有改变意义。但如果还要考虑句式的意义,则

转换会改变意义。在语义解释上,Chomsky 的标准理论体现为短语结构规则生成的范畴和语法关系加上词库项目的插入,Fillmore 格语法体现为深层格关系,生成语义学体现为语义表达式,实际上还有很多其他层面的意义,如表层结构关系(标准理论在短语结构规则中控制的是深层结构关系)、句式语义、语用关系控制的意义,可以说,基础部分只控制了句子某个层面的意义。转换生成语法之所以要在深层结构或句子的逻辑基础上解释语义,是想用简单的语义关系和转换关系来说明复杂的表层结构。

只要是转换,都会不同程度地改变转换式的意义和在语境中的分布。即使那些看上去没有改变意义的转换,如果进行细微的语境分析,总能找出语义差别。比如,当复合动趋式的宾语有数量短语限制的时候,可以有三种位置,这三种转换式的意义看起来似乎是相同的:

	宾语前置	宾语中置	宾语后置
出来	拿一捆书出来	拿出一捆书来	拿出来一捆书
出去	拿一捆书出去	拿出一捆书去	拿出去一捆书
进来	拿一捆书进来	拿进一捆书来	拿进来一捆书
进去	拿一捆书进去	拿进一捆书去	拿进去一捆书
上来	拿一捆书上来	拿上一捆书来	拿上来一捆书
上去	拿一捆书上去	拿上一捆书去	拿上去一捆书
下来	拿一捆书下来	拿下一捆书来	拿下来一捆书
下去	拿一捆书下去	拿下一捆书去	拿下去一捆书
回来	拿一捆书回来	拿回一捆书来	拿回来一捆书
回去	拿一捆书回去	拿回一捆书去	拿回去一捆书
过来	拿一捆书过来	拿过一捆书来	拿过来一捆书
过去	拿一捆书过去	拿过一捆书去	拿过去一捆书

其实这种带宾语的动趋式在句中的分布情况是不一样的。比较：

他从冰箱里拿出来一瓶酒
他从冰箱里拿出一瓶酒来
？他从冰箱里拿一瓶酒出来

我从冰箱里拿出来一瓶酒
我从冰箱里拿出一瓶酒来
？我从冰箱里拿一瓶酒出来

你从冰箱里拿出来一瓶酒
你从冰箱里拿出一瓶酒来
你从冰箱里拿一瓶酒出来

当宾语前置时，在第一人称“我”和第三人称“他”中会受到限制。如果我们继续观察以上句子和“刚才(已经、才、又、昨天)”等已然时间语符以及“正要”(将、要、打算、准备、明天)等未然时间语符的共现情况，就可以看出分布上的差异。先看“刚才(已经、才、又、昨天)”等已然时间语符的共现情况：

他刚才从冰箱里拿出来一瓶酒(后置)
他刚才从冰箱里拿出一瓶酒来(中置)
？他刚才从冰箱里拿一瓶酒出来(前置)

我刚才从冰箱里拿出来一瓶酒(后置)
我刚才从冰箱里拿出一瓶酒来(中置)
？我刚才从冰箱里拿一瓶酒出来(前置)

你刚才从冰箱里拿出来一瓶酒(后置)

你刚才从冰箱里拿出一瓶酒来(中置)

? 你刚才从冰箱里拿一瓶酒出来(前置)

再看“正要”(将、要、打算、准备、明天)等未然时间语符的情况:

? 他正要从冰箱里拿出来一瓶酒(后置)

他正要从冰箱里拿出一瓶酒来(中置)

他正从冰箱里拿一瓶酒出来(前置)

? 我正要从冰箱里拿出来一瓶酒(后置)

我将正从冰箱里拿出一瓶酒来(中置)

我将正从冰箱里拿一瓶酒出来(前置)

? 你正要从冰箱里拿出来一瓶酒(后置)

你正要从冰箱里拿出一瓶酒来(中置)

你正要从冰箱里拿一瓶酒出来(前置)

通过上面的共现,我们可以说:

1. 前置式是表示未然的,所以只能和“正要”等未然语符共现。

2. 中置式既表示未然,也表示已然,所以既可以和“正要”等未然语符共现,也可以和“刚才”等已然语符共现。

3. 后置式只能只表示已然,所以只能和“刚才”等已然语符共现。

前面谈到“你从冰箱里拿一瓶酒出来”可以说,因为第二人称出现在这里可以表示祈使句,祈使句一般都属于未然情况,所以句子成立。带事物宾语的动趋式和已然、未然语符的共现情况可

以概括如下：

	宾语前置	宾语中置	宾语后置
已然	?	+	+
未然	+	+	?

带宾语的动趋式的三种格式在有时体语符的伴随下也能够出现，并且能表达不同的时体，这说明汉语中的时体不仅可以用时体语符来表达，也可以用语序来表达，同时也说明转换或宾语位置的移动改变了时体意义。

现在我们可以说，转换都会改变意义，只不过改变的显著程度不一样。那些看上去没有改变意义的转换式，如果语境扩大，仍然能显示出语义的改变。

现在我们考虑一个更为关键的问题，这是转换生成语法没有讨论的。转换生成语法是用同一个角度对待被动转换、疑问转换、祈使转换、名物化转换的，这个角度就是生成程序。事实上疑问转换和被动转换、祈使转换、名物化转换并不一样。任何NP-VP都可以有疑问转换，但不一定有被动、祈使或名物化转换。比如在汉语中，哪些动词可以有“被动”式，即能否转换成被动，和动词本身的属性有关系。有些动词有“被动”转换，有些没有：

- 他卖了房子→房子被他卖了
- 他通知了大家→*大家被他通知了
- 他缺了帮手→*帮手被他缺了

也就是说，被动转换并不能周遍到所有的动词，但疑问转换可以周遍到所有的动词。比如：

他是不是卖房子？

他是不是通知了大家？

他是不是缺了帮手？

再比如汉语中有一种转换关系：

NP1 + V + NP2	V-NP2-的-NP1
猫追老鼠	追老鼠的猫
男人住一楼	住一楼的男人
小孩姓李	姓李的小孩
学生吃食堂	吃食堂的学生

除了“是”等比较特别的动词，NP1 + V + NP2 总是可以转换成 V-NP2-的-NP1 的。这类转换和被动转换不同，不需要在词库中标注动词的特性，也不需要给动词作进一步的分类。这类转换如果在深层结构中标注，生成句子的时候就可启动相应的转换，不存在选词的困难。

如果要在深层结构中标注“被动”就要知道哪些动词可以有“被动”。从目前对动词的研究情况看，由于被动转换的周遍条件还没有找到，因此能否转换成被动还只能看成动词本身的个性，需要在词库中标注。从某种意义上说，既然在词库中有的动词已经标注了被动，深层结构中是否标注也就没有必要了，因为词项不同深层结构是不同的。不过从句子的生成机制看，仍然需要在深层结构中标注被动。在词库中标注被动显示了动词有转换成被动的性质，在深层结构中标注被动表示要准备生成被动句。

概括地说，从生成的角度看，被动转换要求在词库和深层结构中都有标注，疑问转换只需要在深层结构中标注。从理解的角度看，表层结构的句法标记都在起解释语义的作用。

如果我们将来能够给出平行周遍条件，判定哪些特征的动词有被动转换的可能，我们就可以把动词分成可转换动词和不

可转换动词两类：

- 可被动转换动词 V1：写、砸、拿、修、吃、卖、喝、打、偷……
- 不可被动转换动词 V2：恨、爱、是、信、有、缺……

这个时候是把可转换信息放在词库还是深层结构？我们来做
一个比较：

NP1 + V1 + NP2	NP1-V1-NP2
NP1 + V2 + NP2	NP1-V2-NP2
NP1 + V2 + NP2 (被动转换)	NP2-被-NP1-V2-C (C 是补足成分)

显然，深层结构中 V2 和 V1 的不同类本身就是一种可转换和不可转换的标记。也就是说，在找出动词可以进行被动转换的平行周遍以前，每个动词都要标注是否可以转换成被动式，这个时候可转换标记是在词库中，因此也在深层结构中。如果找出了动词可以进行被动转换的平行周遍条件，就可以通过某个特征提取出一个可以转换的动词小类，这时可以进行被动转换的信息在语类中，因此也是在深层结构中。动词可进行被动转换的特性或语类特征只决定了由该动词构成的深层结构是否有转换成被动的可能，最后被动转换的实现仍然要取决于深层结构中是否已经有了转换的要求，即是否已经有了转换标记。

10.3 平行转换中的语义差别

转换是否改变意义，还和转换的平行条件有关系。

前面分析过，Lakoff 等提出的生成语义学是把量词的差别放在底层的表达式中，即放在句子的深层部分来处理。值得注意的是，Chomsky(1965,3,p136,note9)已经注意到下面现象：

every one in the room knows at least two languages
 at least two languages are known by everyone in the room
 there are two languages that everyone in the room knows

对于任意 NP1 和 NP2, 如果有:

every NP1 knows at least two NP2s

就一定有:

at least two NP2s are known by NP1

这两个句式在量词的约束上有细微差别。这类问题是后来在表层结构解释语义的先兆。

Jackendoff(1969b)最早考虑在表层结构处理语义, 观察下面句子:

(77) (i) Not many arrows hit the target

(没有许多箭射中靶子。否定范围: 名词 many。语义解释: 并非许多箭射中了靶子)

(77) (ii) Many arrows didn't hit the target

(许多箭没有射中靶子。否定范围: 动词短语。语义解释: 许多箭没射中靶子, 但并不否定许多箭射中了靶子)

两句的深层结构标注是相同的:

Not (many arrows hit the target)

但表层意义不完全相同, 这是 not 否定的范围不同。容易看

出, not 否定的范围不是深层结构的内容。否定的范围是由 not 在表层结构的位置决定的。或者说, 在深层结构中标注 Tnot 这样的否定转换, 并不能提供否定范围的信息。这样在转换过程中, 或者说在表层结构中, 必须考虑否定范围问题。Jackendoff 认为像“否定”和“数量”这样一些逻辑成分的控制范围是由表层结构决定的, 除非允许 not 跟构成它的“辖域”(scope)的短语一起出现在深层结构里。否定的“辖域”有高有低, 在“Not many arrows hit the target(并没有许多箭射中)这个句子中, 否定词 not 在句法上是一例量词否定, 作为句子中的最高层算子(它比量词 many 有着更高的辖域), 它还是句子层的否定。然而, 在 Many arrows didn't hit the target(许多箭没有射中目标)这个句子中, 否定词 -n't 是在“小句”层次上对谓词的否定, 因为它的“辖域”要比量词 many 低, 所以它不是句子否定。

观察更多的例子:

(79) the target was not hit by many arrows

(靶子没有被许多箭射中)

(80) not many demonstrators were arrested by the police

(没有很多示威的人被警察逮捕)

(81) many demonstrators were not arrested by the police

(许多示威的人没有被警察逮捕)

(82) John didn't buy many arrows

(约翰没有买下很多支箭)

(83) many arrows were not bought by John

(很多支箭没有被约翰买下)

(84) John bought not many arrows

(约翰买下了不是很多支箭)

(85) not many arrows were bought by John

(不是很多支箭被约翰买下了)

Jackendoff 提出了下面的论点:

如果表层的主语有数量成分,那么句子的否定(如(77i))跟动词短语的否定(如(77ii))在意义上是不同的;但是如果^①数量成分是跟在动词后头的名词短语的一部分,那么句子否定和动词短语否定成分和数量成分的次序就一致了,它们的意义也就相同。

Jackendoff 认为除非在深层结构中标注否定的范围,否则要承认否定的范围是由表层结构决定的。实际上 Jackendoff 是以另一种方式在处理 Lakoff(1971) *On generative semantics* 所提出的量词问题。Lakoff 是要在语义表达式中解决问题,Jackendoff 是要在深层结构处理问题。

这些和语义有关系的因素,比如语用意义,量词、否定等意义,是标准理论没有涉及的。前面提到的生成语义学开始涉及这些因素,并且试图把所有这些意义都放在语义表达式中。由于生成语义学的语义表达式本质上是底层结构,因此从总体上看是把所有这些意义都放在深层结构中。但事实上,这些非语义格层面的意义都是在转换过程中形成的,于是早期 Katz 的转换不改变意义的理论以及 Chomsky(1965)相应的标准理论就出现了矛盾。

Chomsky(1972, p207)特别提到的了 Jackendoff(1969b)的著作 *An Interpretive Theory of Negation*, 承认像否定和数量这样一些逻辑成分的范围是由表层结构决定的,认为表层结构涉及了意义的相同与不同,和标准理论是不一致的,因此标准理

① 根据原文,应该加上“如果”二字。

论需要修改,像量词、否定等意义需要在表层结构中解释。更进一步,Chomsky(1972,p213)认为,由深层结构代表的语法关系决定了语义解释,而焦点、预设、话题、评价、指涉、逻辑因素(量词、否定)等则是表层结构决定的。我们的问题是,焦点、预设、话题、评价、指涉、逻辑因素(量词、否定)这样一些转换和被动、疑问、名物化等转换有什么根本区别,为什么后者要在深层结构解释语义?

我们认为这仍然是“转换不改变意义”这说法的含混性带来的麻烦。如果我们绕过“转换不改变意义”的字面含义,则可以更深入认识转换和语义的关系。

其实,如果允许在底层标注表层结构的区别,量词、否定词的辖域仍然可以像被动、疑问等一样在深层结构中标注。在深层结构解释语义还是在表层结构解释语义,这仍然只是技术问题,在实际语言的识别中,人们仍然是从量词和否定词出现的表层位置来理解量词、否定词的辖域的。因此,从严格意义上说,转换引起的意义改变都不是深层结构蕴含的。

我们认为,转换和语义更重要的关系还在于转换会改变同一语义结构的不同句式在语义上的平行性,这是生成语法考虑较少的。而弄清这个问题对自然语言的理解条件和生成条件有重要意义。现在我们先从平行关系的角度来认识转换改变意义或平行关系的实质。一般地说,转换往往能保持一种平行关系。前面已经提到下面的转换是平行的:

John broke the cup→The cup was broken by John

John ate the apple→The apple was eaten by John

下面的转换也是平行的:

Many men read few books→Few books are read by

many men

Many men bought few books → Few books were
bought by many men

如果所有的主动句和被动句都能保持以上平行关系,那么被动句的深层结构只需要在主动句的深层结构上标注被动转换 Tpassive,被动句的意义就可以在被动深层结构上得到解释,从被动句深层结构向被动句表层结构转换就不再改变意义,或者说不再处理意义,这时转换只是一个技术操作问题。

但是,如果把上述两组实例进行比较,情况就很不一样:

1. John broke the cup → The cup was broken by John

2. John ate the apple → The apple was eaten by John

1. Many men read few books → Few books are read by
many men.

2. Many men bought few books → Few books were bought
by many men

这时左右两端的关系不平行了。我们认为所谓改变意义的更关键问题就是在这里。不是说主动和被动意思没有变化,而是在转换的过程中,一旦涉及量词(或否定),主动和被动式之间的转换关系就不平行了,或者说主动和被动之间产生的语义变化就不一样了,这时候在深层结构上标注被动转换和解释语义就出现了问题。

现在归纳一下我们对语义解释层面的看法。语义结构意义、词汇意义都是深层结构语义,转换造成的句式意义都是表层语义,即不仅数量、否定、焦点等转换是表层意义,被动、疑问、名物化转换等也都是表层意义。

如果我们从句子生成和句子理解的过程看,句式的意义跟深层结构和表层结构都有关系。从句子的生成过程看,如果在深层结构标注了被动,就要求被动转换来形成句子的表层结构。从句子的理解过程看,表层的被动标记决定了句子的被动含义,主动标记决定了句子的主动含义。一般地说,转换带来的位置、标记等的变化不仅是判定表层语义的关键,也是判定语义格的关键。后来 Chomsky(1981)和其他一些学者全部在表层结构解释语义,我们认为从语言理解的角度看是有道理的。但从生成句子的角度看,情况不完全一样。

思考问题

1. 转换的平行周遍条件应该从哪些方面去寻找?
2. 转换和语义的关系是什么?

11. 结构的阶与递归性

前面说从 Bloomfield 开始,结构语言学尽可能不使用句法结构关系或句法功能的概念,不过 Bloomfield 使用了向心结构和离心结构两个概念。从一个角度看,Bloomfield 是想根据分布功能来定义结构关系,所以向心结构只涉及语类是否相同的问题,并没有直接提到结构关系。从另一个角度看,由于向心结构有一个核心,于是另一个成分自然就成了补足语,这就在一定程度上涉及了偏正结构的问题。另外,Bloomfield 提到了支配的概念,这也和结构关系有联系。

前面讨论过,Harris(1946)彻底放弃结构关系,试图从语类推导出结构关系。Hockett(1954)试图说结构关系是初始概念,实际上不成功,因为他给出的实例的结构关系可以从整体的语类加以定义,实际上仍然可以通过语类推导出来,只不过不是从直接成分的语类推导出来,而是从整体的语类推导出来(陈保亚,1999,2.1.3)。Chomsky(1965)用语类和语序来定义句法结构关系和语义格,等于是取消了句法结构关系和语义结构关系的独立性,统一起来说是取消了结构关系的独立性。功能、结构关系、句子成分,都是说的同一个问题。从 Harris(1946)开始,结构语言学力图取消功能、结构关系、句子成分等概念,Chomsky(1957,1965)的基本思路也是这样,这就形成了后来的形式语法和功能语法的对立。

下面我们从 X 杠理论的分析中进一步认识结构关系的重要性。

11.1 X 杠理论

Wells(1947)已经注意到名词短语有两类:

king of England

the king of England

对于汉语来说,也有类似的情况:

英国国王

这位英国国王

在大多数语言中,都存在这两类名词短语,区别在于是否有限定成分(determiner)。这两类名词短语在分布上是不同的。Radford(1988,4.3)归纳了一些经验证据:

The king of England opened Parliament

* *king of England* opened Parliament

They crowned the king of England yesterday

They crowned * *king of England* yesterday

Parliament grants little power to *the king of England*

Parliament grants little power to * *king of England*

He became [king of England]

the [English king]

the [King of England]

Who would have dared defy the [King of England] and
[ruler of the Empire]

(Ordinary Coordination)

He was the last (and some people say the best) [king of
England]

(Shared Constituted Coordination)

The present [king of England] is more popular than the
last *one*

*The [king] of England defeated the *one* of Spain

(Pronominalization)

在 X 杠理论出现以前, 名词短语的描写是这样的:

$$NP \rightarrow D \ N$$

(Chomsky1957: $NP \rightarrow T + N$; Chomsky1965: $NP \rightarrow Det \cap N$)

但是, 这样的改写规则不能区分上面提到的无定和有定两类名词短语。要描写上面的两类名词短语的区别, 可以考虑这样的一套规则(以英语为例子):

$$NP \rightarrow D \ NP$$

$$NP \rightarrow N \ PP$$

根据这套规则可以生成下面的短语:

$$\langle {}_{NP}({}_D \text{ the}) ({}_{{}_{NP}}[{}_N \text{ king}] [{}_{PP} \text{ of England}]) \rangle$$

这套规则尽管能够生成出上面两类名词短语,但第一个规则“NP→D NP”反复使用会生成下面结果(不断用“D NP”改写NP):

$$\begin{aligned} \text{NP} &\rightarrow \text{D NP} \\ &\rightarrow \text{D} [\text{D NP}] \\ &\rightarrow \text{D} (\text{D} [\text{D NP}]) \end{aligned}$$

比如:

- * a our room
- * the our room
- * the the room

两类名词短语的性质也没有在规则中体现出来。不仅名词短语存在层次上的区别,动词短语、形容词短语等也存在类似的区别,也都有改写规则所遇到的问题。

Chomsky(1970, *Remarks on Nominalization*)提出了一种语杠理论,或者说 X 杠理论。根据该文的改写规则(48)和(49),我们对符号略加以调整,用数字表示语杠,其基本模式可以概括为:

$$\begin{aligned} \text{X}_2 &\rightarrow \text{spec X}_1 \\ \text{X}_1 &\rightarrow \text{X Comp} \end{aligned}$$

对于 the king of England,用 Chomsky 语杠理论可以分析为:

$$\langle \text{X}_2 (\text{spec the}) (\text{X}_1 [\text{X king}] [\text{Comp of England}]) \rangle$$

其中 Spec 是限制成分,Comp 是补足语成分,X 右下方的数字表

示 X 的杠。X 杠理论试图概括短语结构的一般性质,即任何短语都有一个核心成分 X, X 和 Comp 一起构成 X1 的直接成分,限定语 Spec 和 X1 构成 X2 的直接成分, X2 是核心成分 X 的最大投射。

语杠理论是对名词短语、动词短语、形容词短语、介词短语、副词短语以及句子的统一表述。不过句子的语杠分析在生成语法中比较特殊。比如:

He was waiting for the king of England

其深层结构层面的语杠分析是:

{_{IP} He <_I [_I INFL] [_{VP} be waiting for the king of England]>>}

句子被分析成屈折短语(IP), 屈折成分(时态、人称、数)是 IP 的核心, VP 是屈折成分 I(INFL, inflection)的补足语, I 和 VP 共同构成 I', 主语 He 是 I' 的限定语 Spec。而在 VP 中, 核心是动词, 限定语 Spec 是 be。如果是嵌套句, 还有 CP 分析。比如:

John knew {_{CP} <_C [_C that] [_{IP} Jack was waiting for the king of England]>>}

即 knew 的补足句或补足短语 CP(Complementiser Phrase)的核心是 that, 这里的 that 又叫 Complementiser 或者 Comp, 子句 IP 是 that 的补足语。

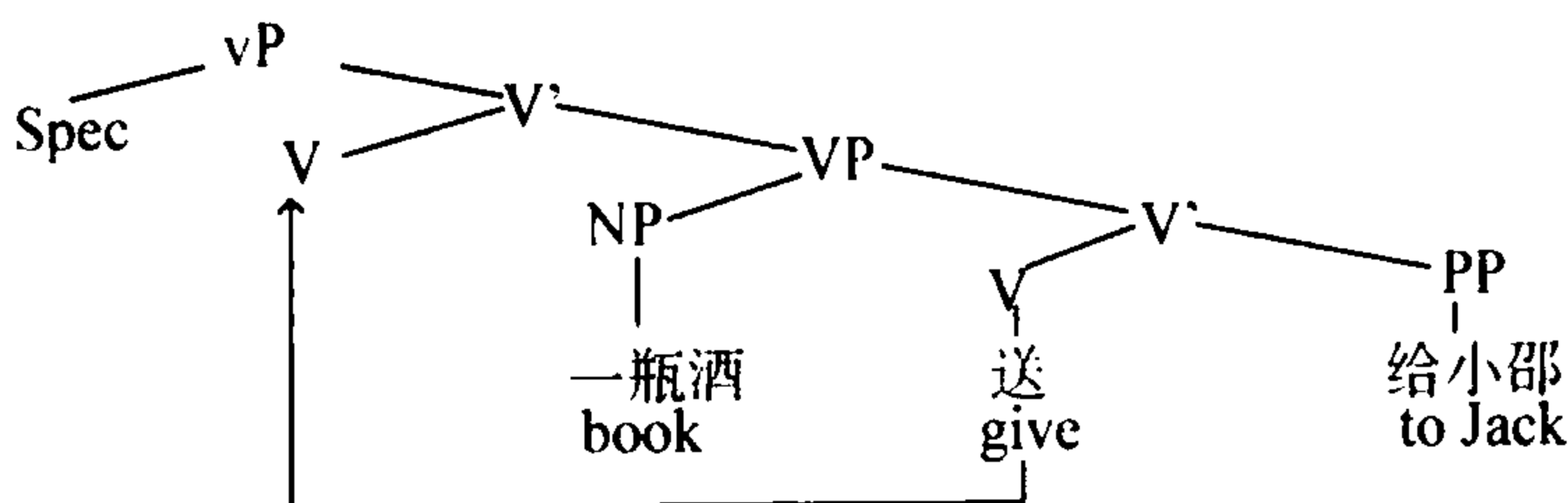
IP、CP 这种分析在英语中是有一定道理的, 我们认为其依据在于英语的句子是定式动词。但汉语的情况不完全相同。朱德熙(1985)提出过词组本位的思想, 认为汉语词组结构和句子结构在形式上没有根本区别, 可以用同一种分析方法处理, 也有

一定的根据。在后面的分析中,如果不影响讨论,我们一般不严格区分 VP 和 IP。

语杠理论通常采取的是结构二分法(binary branching),即一个上位成分被分成两个下位成分。但自然语言中确实存在一些可能多分的情况,比如:

[送] [一瓶酒] [给小邵]
[give] [a book] [to Jack]

为了在理论上进行统一, Larson(1988)提出了轻动词 *v* (light verb)的概念解决以上双宾语结构的二分问题。上面的实例用轻动词分析如下:



这里的要点是区分了 vP 和 VP。最高层 vP 的核心是轻动词 *v*, 第三层 VP 的核心才是动词“送”,因此在深层结构中,双宾语是一个 vP 统辖一个 VP,这里的 VP 被称为 VP-shell 或 shell structure。通过转换,“送”移动到轻动词的位置,才形成双宾语结构。英语的情况类似。

语杠理论的思想最早始于 Zelling Harris(1946, 1951)。Harris 在谈到替换的时候已经注意到可替换性的条件。比如:

DA=A: quite old = old
N-s=N: paper + -s = paper

如果用 DA 替换 A, 比如用 quite old 替换 realy old 中的 old, 就有:

$$DDA = A: \quad \text{realy quite old}$$

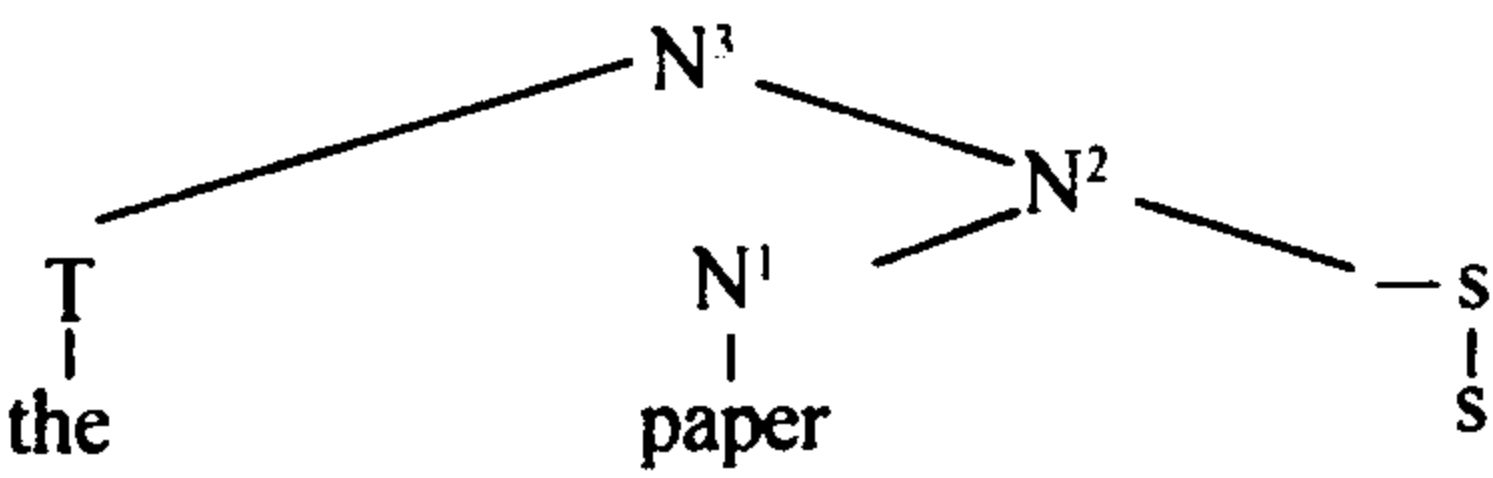
但是, 如果用 N-s 替换 N, 比如用 papers 替换 paper, 尽管 pa- pers 和 paper 的语类是相同的, 但 papers 替换 paper 会出现下面现象:

$$\text{papers} \rightarrow ^* \text{paperss}$$

Harris 注意到的这个现象并不是一个符号技巧的问题, 而是涉及替换本质的问题。当时 Harris 是用上标数字表示名词短语在替换性质上的区别:

$$\begin{aligned} N^1-s &= N^2 \\ TN^2 &= N^3 \end{aligned}$$

根据 Harris 的思路, 我们可以得到以下替换结构:



这已经具备了 X 杠理论的层阶思路, 但还没有核心和附加的概念。和 Harris(1946) 的标记方式不同, Chomsky(1970) 以及后来的分析是用上加横线表示名词短语的区别。Jackendoff (1977a, X-bar Syntax: A Study of Phrase Structure) 用 prime notation 表示名词短语的区别。比较:

Number notation	bar notation	prime notation
N^1 (N-one)	N	N
N^2 (N-two)	$\bar{N}$ (N-bar)	N' (N-prime)
N^3 (N-three)	$\bar{\bar{N}}$ (N-double-bar)	N'' (N-double-prime)

尽管所用的符号不同,但比较三者的思想可以看出,核心问题是名词短语的分布差异。为了印刷的方便,我们后面的讨论将用 Harris 的数字标记法来区分名词以及其他词的杠。为了在线性分析中体现杠的层次,不引起混乱,我们给出一种直接成分同标法,即互为直接成分或相互控制的成分用相同的括号等级。于是有:

$\langle N_2 (_{Spec} \text{ 这张}) (N_1 [_{Comp} \text{ 小}] [N \text{ 桌子}]) \rangle$

“小”和“桌子”互为直接成分,都用方括号[],“这张”和“小桌子”互为直接成分,都用括号()。更多的层次以此类推。

一般认为,X 杠理论的重要意义是对名词短语、动词短语、形容词短语、介词短语等进行统一描述,因为各类短语都有核心、Comp 和 Spec。实际上这只是次要的方面,因为结构语言学描写直接成分就是超越具体语类,把直接成分作为各种语类短语的共同的本质加以描述的。

和 Chomsky(1957)短语结构规则相比,语杠理论首先是对短语结构规则的简化,这是 Chomsky(1965)次范畴化理论提出以后的必然趋势,因为 Chomsky 当时的次范畴化是在短语结构规则和词库中同时进行的。比较汉语的情况:

短语结构规则	词库中词项 的次范畴	实例
VP→V		
VP→V+NP	学习_NP	学习英语
VP→V+NP1 +NP2	给_NP1_NP2	给小邵英语书
VP→V+VP	喜欢_VP	喜欢学习英语
VP→V+S	知道_S	知道汪锋喜欢学习英语
VP→V+NP +S	告诉_NP_S	告诉小邵汪锋喜欢学习英语
.....		

前面我们讨论过，词库中的次范畴化是必须要做的工作，因为一个动词能带哪些必要论元是找不到平行周遍条件的。既然词库中的词项已经次范畴化，再在短语结构规则中次范畴化，工作就重复了，所以短语结构规则部分更应该研究直接成分的一般性质。

我们认为 X 杠理论最重要的价值是对各类短语结构有了更深的认识，除了语类和直接成分，还包括最大投射。结构语言学的直接成分理论认为结构的组合是有层次的，但结构语言学的层次主要是从表层结构入手，而语杠理论是从深层结构的角度认识直接成分的性质，尤其是在传统语法各种纷繁的句子成分关系中认识到核心和附加的重要性。尽管后来 Chomsky 本人以及其他学者对语杠理论有修正，有的甚至放弃语杠理论，但我们认为这只是形式化描写的技术问题，语杠理论所涉及的自然语言的一些根本性质，仍然是根本的。要认识这一点，再回到 Wells 的例子上来：

the king of England

下面两种层次似乎都有一定的道理:

1. the [king of England]
2. the king [of England]

按照 X 杠理论,只有 1 是恰当的。X 杠理论是把各种语类统一起来论述的,为了便于理解,下面用具体语类来标注 1 的结构:

$\langle_{N^2} \text{ the } (\langle_{N^1} [\text{king}] [\text{pp of England}]) \rangle$

从以上语类标注可以看出,X 杠理论的一个核心思想是对结构中直接成分的认识。以上短语包括了以下一些必要部分:

N^2 : the king of England

N^1 : king of England

N : king

PP : of England

可以概括为以下几个方面:

1. N^2 是名词短语的最大投射,或者说是结构上充分的名词短语
2. N^1 是包含在 N^2 中间的短语,是结构不完备的名词短语
3. N 是 N^2 的核心
4. pp 是 N 的补足语
5. the 是 N^1 的限定语

名词的最大投射是指以名词为中心形成的一个完整的名词短语 (full nominal phrase) 所需要的必要成分。关于“完整”的概念后面我们会加以解释。以上的 the、of England 都是 N^2 的必要成分。由于 X 杠理论在短语中引入了最大投射的概念, 对直接成分和层次的理解就和结构语言学有差异。通过和结构语言学直接成分的比较, 可以看出这种区别:

$$\langle_{N^2} \text{ the } (\langle_{N^1} [\text{ }_N \text{ king}]) \rangle$$

按照结构语言学的直接成分理论, the king 中的 the 和 king 是直接成分关系, 即姊妹关系, 但按照 X 杠理论, the 和 N^1 是姊妹关系, 和 N 不是姊妹关系, N 是 N^1 统领 (dominate) 的成分, 因此 the 和 king 是间接成分的关系, 或者说, king 是 the 的侄子。比较前面的树形图可以看得更清楚。X 杠理论实际上是在从深层结构的角度理解层次或直接成分关系, 在后面我们会看到, 有一类层次引起的歧义现象只能用 X 杠理论能够解释, 因此它的存在是必要的。

按照以上的 X 杠理论, 下面汉语的名词短语可以分析成:

$$\langle_{N^2} (\text{Spec 这张}) (\langle_{N^1} [\text{Comp 小}] [\text{ }_N \text{ 桌子}]) \rangle$$

$$\langle_{N^2} (\text{Spec 这张}) (\langle_{N^1} [\text{ }_N \text{ 桌子}]) \rangle$$

在后一句中, spec(这张) 和 N^1 是直接成分关系, spec 和 N (桌子) 是间接成分关系。这和通常从表层结构理解直接成分关系是不同的。再考虑更大的语境:

$$1. [\text{ }_V \text{ 买}] / \langle_{N^2} (\langle_{N^1} [\text{ }_N \text{ 桌子}]) \rangle$$

$$2. [\text{ }_V \text{ 买}] / \langle_{N^2} (\langle_{N^1} [\text{Comp 小}] [\text{ }_N \text{ 桌子}]) \rangle$$

3. $[_V \text{买}]/\langle N^2(\text{Spec 这张})(N^1[_N \text{桌子}])\rangle$
 4. $[_V \text{买}]/\langle N^2(\text{Spec 这张})(N^1[_{Comp} \text{小}][_N \text{桌子}])\rangle$

从结构语言学的表层结构看,斜线“/”两端的词或组合相互之间是直接成分的关系。按照 X 杠理论,只有 4 中“买”和“这张小桌子”是直接成分关系。不过“直接成分”这个术语是结构语言学中的,不能准确区分深层结构和表层结构中的直接成分关系。为了更准确区分这种关系,我们可以通过家庭亲属关系的比喻引入同辈直接成分、不同辈直接成分、同辈姊妹关系和不同辈姊妹关系几个概念。同时引入“空成分”和“完备”的概念,带有空成分的短语是不完备的。这时我们就可以说 1、2、3 例中 N^2 的直接成分是不同的辈直接成分,因为 N^2 中存在空成分。4 中 N^2 的直接成分是同辈直接成分。我们说 1 中的“桌子”是“买”的不同辈姊妹成分,因为 N^2 中的限定语(Spec)和补足语(Comp)都是空的。由于这时的“桌子”缺少限定语和补足语,我们说这时的“桌子”是缺杠的,具体说是缺 2 杠的。同理,2、3 中 N^2 的两个直接成分的关系也是不同辈姊妹关系,只不过是缺 1 杠的。2 中的“小桌子”是“买”的不同辈直接成分,因为这时 N^2 中缺少限定语(Spec)。由于这时的“小桌子”缺少限定语,我们说这时的“桌子”也是缺杠的,具体说是缺 1 杠的。3 中的“这张桌子”是“买”的不同辈直接成分,因为这时 N^2 中缺少补足语(Comp)。由于这时的“这张桌子”缺少限定语,我们说这时的“桌子”也是缺 1 杠的。4 中 N^2 的两个直接成分的关系是同辈直接成分关系。

正是有 1、2、3 实例中的 N^2 有不完备的直接成分,有空成分存在,所以这时的 N^2 是可以进行扩展替换的。我们认为这是 X 杠理论的另一个重要价值之一。至于完备直接成分是否可以被扩展替换,这需要一定的条件,后面会展开。

X 杠理论和结构语言学直接成分理论的另一个区别是重新

引入核心的概念。Bloomfield(1926, 1933)并不怎么讲句子成分的问题, 他的句法结构不是从主谓、述宾等角度定义的, 而是从直接成分的功能和整个结构的功能的异同来定义的, 这就是著名的“向心结构(endocentric construction)”和“离心结构(exocentric construction)”的概念, 短语的“核心”成为重要概念。向心结构指至少有一个直接成分的功能是和结构的功能相当的结构, 和结构的功能相同的直接成分是核心(head)。离心结构指没有一个直接成分的功能和结构的功能相当的结构。Bloomfield 对于向心结构的定义如下(这里的功能即语素):

合成短语(resultant phrase)可能和一个(或多个)成分一样属于同一个形类。譬如, poor John 是一个专有名词短语, John 这个成分也是这样; John 和 poor John 这两个形式, 从整体看, 具有同样的功能。因此, 我们说英语的特性—实体(character-substance)的结构(如 poor John, fresh milk, 等等)是一个向心(endocentric)结构。(Bloomfield, 1933, p194, 中译本 1997, p239)

Hockett(1954)曾经分析过三种不同的语言分析模式, 尤其是其中的项目与配列模式(IA)和项目与变化模式(IP)最有代表性。我们认为这两种模式最根本的区别在于 IP 模式强调基式和派生, 强调成分之间的句法结构关系, 而 IA 模式主张用分布来取代这两种观念。上面 Bloomfield 关于向心结构的定义实际上还保留有 IP 模式的部分分析思想, 即跟结构关系有联系的核心和附加的关系。后来 Harris(1946, 1951)、Wells(1947)更加强调 IA 分析模式, 在短语中只考虑语类和层次, 不再使用核心的概念, 目的显然是要更彻底地放弃结构关系的概念。Chomsky(1965)讨论过不用结构关系的理由, 认为结构关系已经蕴含在语类中。X 杠理论重新引入核心, 并认为每种短语都

有核心,实际上是重新引入了结构关系,不过只承认偏正关系,并没有引入主谓、述宾等关系,具体的各种结构关系仍然留给语类解释。X 杠理论中关于核心的概念也和 Bloomfield 的核心不同,所有的结构都是有核心的,可以是向心结构的核心,也可以是离心结构中的虚词,甚至包括形态。

如果我们从直接成分的角度来分析语杠的两杠理论,容易看出 X 两杠理论已经包含层次的内容,因为把以上两个规则各使用一次就有:

$$\text{Spec [X Comp]}$$

假设把 X 两杠理论改写成以下一杠模式:

$$X^1 \rightarrow \text{Spec X Comp}$$

虽然也能够生成同样的序列,但层次属性不能体现出来,所以 X 一杠理论描写短语是不充分的,至少要有二杠。

以上 X^2 杠理论是限定语在前,补足语在后。像汉语这样的语言,名词短语的边际成分都是在前,是否也一定需要 X^2 杠短语来描写? 后面会讨论这个问题,并证明至少需要 X^2 杠理论。

11.2 限定语和补足语层次的经验依据

从上面的实例看,一个完整的 $XP(X^2)$ 必须要有限定成分和补足语,限定成分受 X^2 直接统辖(dominate)的,是 X^2 的第一代直接成分,补足语是受 X^1 统辖的,是 X^1 的第一代直接成分。现在考虑:

[the] [king of England]

[the king] [of England]

这种层次差异的证据是什么？下面仍然以 NP 为例来论证。

确定 X 杠分支所遇到的问题也是早期直接成分理论在确定直接成分时遇到的问题。Wells 在他的直接成分中采取的是第一种方式 (Wells 1947), X 杠理论和 Wells 的分析在这一点上是一样的。Wells 当时依据的理由是结构对称 (structure symmetry)。比如：

(a) the [English king]

(b) the [king of England]

当时 Wells 并没有讨论 king of England 的性质, 也没有命名。前面已经提到, king of England 和 the king of England 两个片断的分布是不相同的。我们还可以通过范继淹的并立扩展法进一步论证 Wells 上述层次切分的合理性。范继淹 (1963) 提出了并列扩展, 如果 ABC 的后两项能扩展为并立结构 ($BC + B'C'$), 则 $ABC = A(B + C)$ 。例如“他很高”, 可扩展为“他很高, 很瘦”, 所以层次组合是“他 | 很高”。反之, 如果 ABC 组合的前两项能扩展为并立结构 ($AB + A'B'$), 则 $ABC = (A + B)C$ 。例如“红的花”可以扩展为“红的和白的花”, 所以层次组合是“红的 | 花”

我们曾经对并列扩展法的合理性做过讨论 (陈保亚, 1999, 2.1.4), 可以给范继淹的并立扩展法一种更形式化的描述:

如果 $ABC \rightarrow A([BC][B'C'])$, 其中 B 和 B' 可以相同, 或者 C 和 C' 可以相同

则 $ABC = A(BC)$

如果 $ABC \rightarrow ([AB][A'B'])C$, 其中 B 和 B' 可以相同, 或者 A 和 A' 可以相同
 则 $ABC = (AB)C$

根据范继淹的方法, 上面 Wells 分析英语 the king of England 的层次可以得到确认, 应该是 (the)(king of England)。下面来具体分析。

Radford(1988, 4.3)给出了好几个原则来证明 the king of England 的层次, 其中关键的原则是并接原则(Ordinary coordination), 比如:

Who would have dared defy the [King of England] and
 [ruler of the Empire]
 (Ordinary Coordination)

其实这一原则正好和范继淹(1963)并列扩展法相似。令:

A = the B = King C = of England

于是 the king of England 可以作以下并列扩展:

$[_A \text{ the }] [_B \text{ King }] [_C \text{ of England }]$
 $\rightarrow \{ [_A \text{ the }] \} \{ ([_B \text{ King }] [_C \text{ of England }]) \text{ and } ([_B \text{ ruler }] [_C \text{ of the Empire }]) \}$

不过, 并列扩展法属于一种表层分析。比较:

小王扔了桌子 →

小王扔了桌子、摔了杯子

* 小王扔了、民工捡了桌子
桌子小王扔了, 小李捡了。

第二个扩展不成立, 但如果桌子在表层主语位置, 扩展就成立, 即第三个例子。

不过我们还可以从功能上来论证限定语和补足语在不同层次上。限定成分的本质是使 NP 获得指称, 而补足语从本质上说是一种分类。如果 NP 不含限定成分, 修饰语再多也没有指称:

students of Peking University
excellent students of Peking University

汉语的情况类似:

北京大学学生
北京大学优秀学生

当限定成分加上以后, NP 就有了的确定的或不确定的指称:

a student of Peking University
three students of Peking University
the students of Peking University

an excellent students of Peking University
three excellent students of Peking University
the excellent students of Peking University

一位北京大学学生
三位北京大学学生
那些北京大学学生
一位北京大学优秀学生
三位北京大学优秀学生
那些北京大学优秀学生

一旦加上限定成分,短语就不是表示分类而是指称,NP 的外层就不能再加修饰成分了,因此限定语是比补足语层次更高的成分。

所谓 NP 含有限定语后不能再加修饰,从另一个角度理解就是已经满足最大投射。结构的这种不能再加修饰的概念,是结构的重要属性。Bloomfield(1933,12.11)已经谈到 closed 的概念:

If, however, we join a form like *milk* or *fresh milk* with the attribute *this*, the resultant phrase, *this milk* or *this fresh milk* has not quite the same function as the head or center, since the resultant phrase cannot be joined with attributes like *good*, *sweet*: the construction in *this milk*, *this fresh milk* is *partially closed*. The possibilities in this direction, in fact, are limited to adding the attribute *all*, as in *all this milk* or *all this fresh milk*. When the attribute *all* has been added, the construction is *closed*; no more attributes of this type (adjectives) can be added.

即当 *milk* 或 *fresh milk* 前面没有 *this* 时,可以添加很多成分:

milk→sweet milk→good sweet milk

fresh milk→sweet fresh milk→good sweet fresh milk

一旦带上 this 以后,前面就不能再加除了 all 以外的成分了:

milk→this milk→all this milk→ * sweet this milk

fresh milk→this fresh milk→ * sweet this fresh milk

以上我们主要从英语来论证限定语(Spec)和补足语(Comp)。由于英语的名词短语中,限定语(Spec)和补足语(Comp)分别位于核心两端,所以要追问这两类修饰成分的层次。对于有些语言,比如汉语,名词短语中限定语和补足语都在中心语的左边,所以它们的层次是很清楚的:

$\langle N^2 (\text{Spec 这位}) (N^1 [\text{Comp 英国}] [N \text{ 国王}]) \rangle$

这里实际上依据了一个临近原则,即依照临近核心词的远近来确定层次。我们认为汉语这种语言为 X 杠理论中限定语和补足语的层次提供了进一步的证据,即限定语(Spec)直接受 N^2 统辖(dominate),和 N^1 互为直接成分,而补足语(Comp)直接受 N^1 统辖,和 N 互为直接成分。更概括地说,限定语(Spec)直接受 X^2 统辖(dominate),和 X^1 互为同辈分姊妹成分,而补足语(Comp)直接受 X^1 统辖,和 X 互为同辈分姊妹成分。

11.3 附加成分的再分类

后来的研究证明,附加成分还可以进一步分成 complement 和 adjunct,这种区别已经暗含在 Chomsky(1965, p101-103)的次范畴化思想中。

根据下面的 X 杠理论：

$$X^2 \rightarrow \text{Spec } X^1$$

$$X^1 \rightarrow X \text{ Comp}$$

只能分出限定语和补足语。Chomsky(1965, p101)注意到了下面实例有歧义：

he decided on the boat

这个句子既可以理解成“他决定选择船”，也可以理解成“他在船上作决定”。Chomsky 所举的以上实例，在结构语言学的表层结构层次分析中层次是一样的：

he (decided [on the boat])

可见结构语言学的直接成分分析或表层的层次分析还不能区分出这里的歧义。转换生成语法 50 年代的短语结构规则也不能区分这里的歧义。引起这种歧义根源就在于有两种不同的 PP，它们处在不同的层次上。Chomsky(1965, p103)就两类 PP 还给出了一些实例：

A 类：

He argued [with John][about politics]

He aimed the gun [at John]

He talked [about Greece]

He ran [after John]

He decided [on a new course of action]

B 类:

He died [in England]

John played Othello [in England]

John always runs [after dinner]

Chomsky 认为前一组方括号中的成分能起到次范畴化 (sub-categorization) 的作用, 而后一组不能。如果我们从规则和不规则的角度看, 前一组是固定组合, 介词的选择也是固定的, 后一组是规则组合。后来有的学者把 A 类叫做补足语 (Comp), B 类叫做附加语 (Adjunct)。通过这种区分, 在表层结构中直接成分相同的线性组合, 在深层结构中层次和树形图并不相同:

[_V decided] [_{Comp} on the boat]

[_V decided] [_{Adjunct} on the boat]

[_{V'} [_N decided] [_{Comp} on the boat]] [_{Adjunct} on the boat]

上一句的 PP 是补足语 (complement), 下一句的 PP 是附加语 (adjunct), 第三句包括两种情况。以上三种 V 阶的格式树形图分别是 (见图 11.3.1, 图 11.3.2):

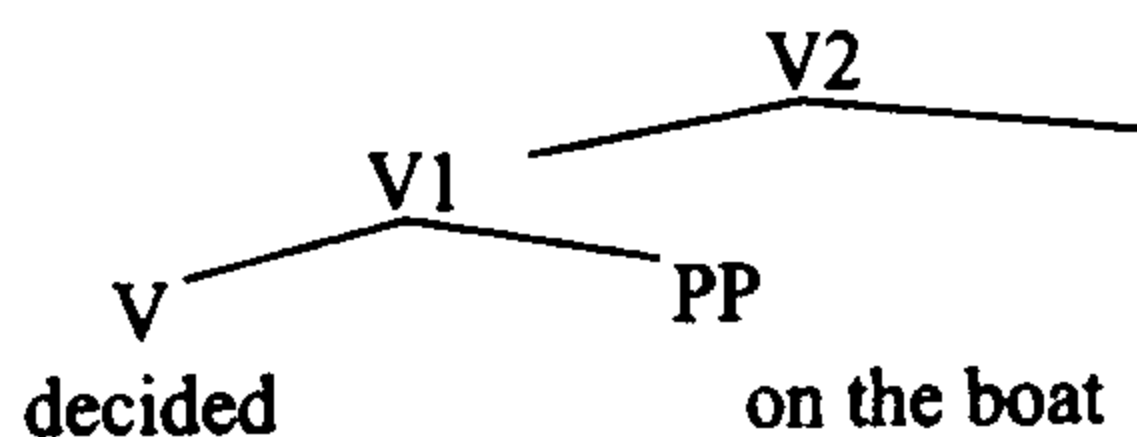


图 11.3.1 补足语动词短语 V 阶格式的树形图

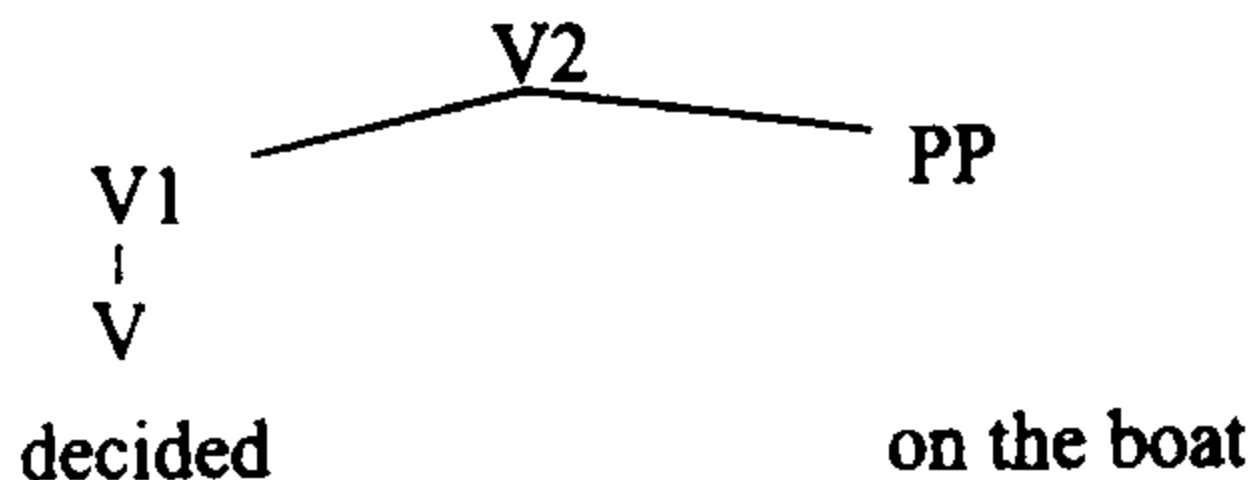


图 11.3.2 附加语动词短语 V 阶的格式树形图

既有附加语又有补足语的短语是(见图 11.3.3):

he { v' (v [v decided] [$_{pp}$ on the boat] v') ($_{pp}$ on the boat $_{pp}$)}

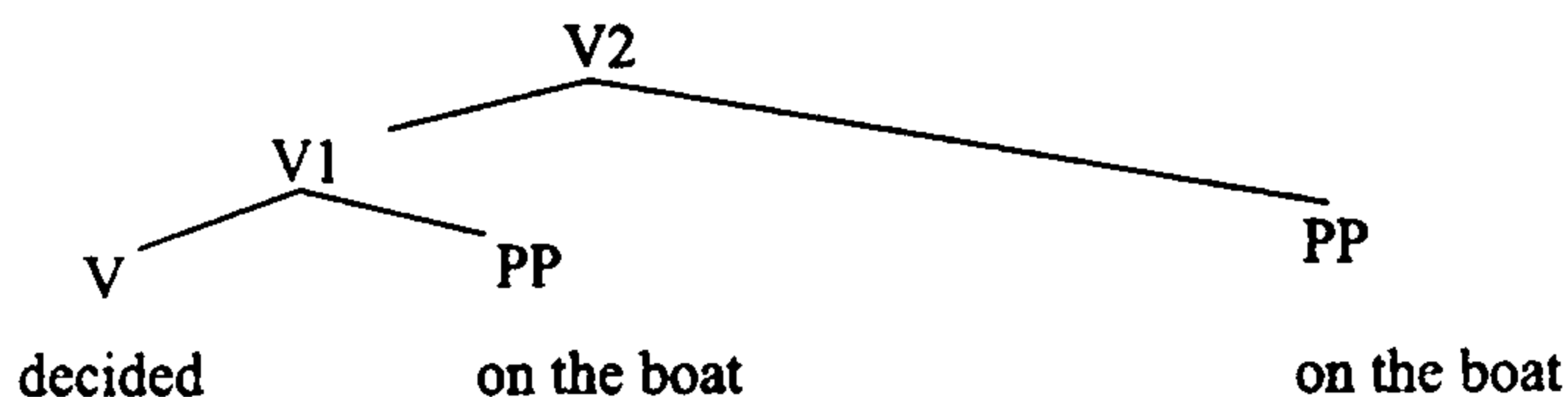


图 11.3.3 补足语附加语动词短语 V 阶格式的树形图

再回到 Chomsky 所列举的动词修饰语的区别上来。如果我们从扩展替换的角度看,在 A 类实例中,VP 的核心 V 不能用 B 类的 VP 进行扩展替换:

He < (argued [with John]) (about politics) >
 → * He < ([argued in classroom] [with John]) (about politics) >

类似的例子还有:

- * He [aimed in classroom] the gun [at John]
- * He [talked in classroom] [about Greece]
- * He [ran yesterday] [after John]
- * He [decided yesterday] [on a new course of action]

相反,在 B 类实例中,VP 中的核心 V 可以用 A 类中的 VP 替换:

He decided in the classroom

He [decided on the boat] in the classroom

从另一个角度看,以上分析把补语的差别转化成了组合层次上的区别,这种转化的合理性下面还要讨论。

以上是关于动词附加成分的分析。Chomsky(1965)所观察到的动词的修饰成分的差异,我们也可以转换成名词短语来观察:

the decision on the boat with my friends

Jackendoff(1977a,p58)在名词短语中也区分了两种不同的修饰语,Hornstein 和 Lightfoot (1981)也观察到名词的修饰语也同样存在差异。比如:

the student [of Physics][with long hair]

* the student[with long hair][of Physics]

of Physics 可以成为 complement,而 with long hair 可以成为 adjunct。一般地说,complement 总是在 adjunct 前面,这仍然符合前面讨论的直接成分临近原则。

临近原则是说离核心越紧密的附加成分,在线性位置上距离核心越近。一般地说,当一个核心带有两个附加成分时,可以有下面三种位置:

[小][红]房子

送[他][一本书]

[经常]读[唐诗]

前两种情况因为附加成分都集中在核心的前面或后面,所以用距离原则就能判定哪个成分距离核心最近。但第 3 种情况由于附件成分分别分布在核心的前后,就有一个判定哪个附加成分

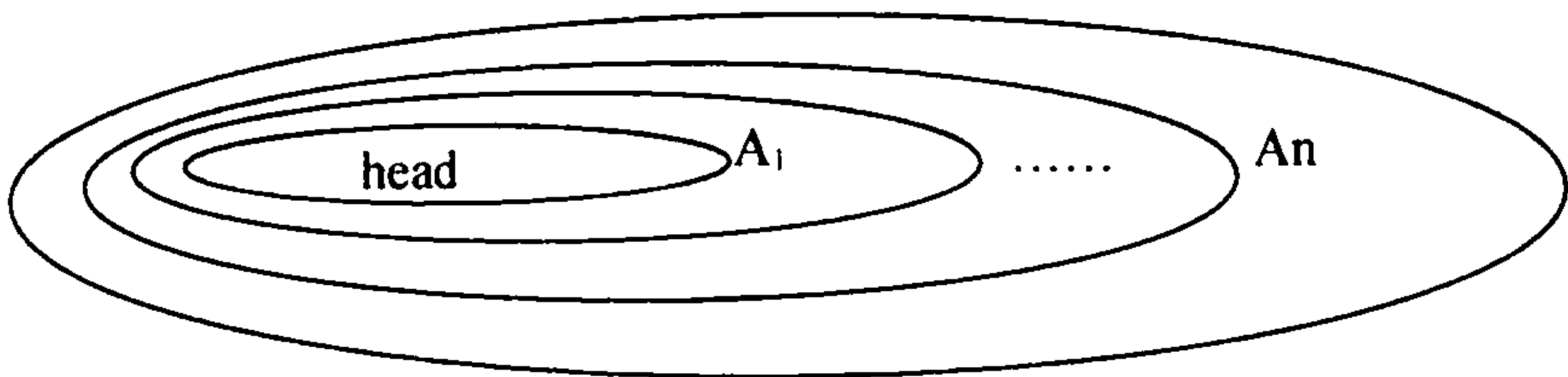
更接近核心的问题。当附加成分超过两个的时候,并且在核心前后都有分布的时候,都存在判定问题:

- [别]读[唐诗][了]
- [别][再]读[唐诗][了]
- [刚刚]读[完][了][一首唐诗]

考虑到不连续直接成分,情况更复杂:

刚刚毕了业

这说明层次的根本问题不是线性问题,而是附加成分距离核心远近的问题。如果一个核心(head)有 n 个附加成分 A,就需要判定哪个是 A_1 ,哪个是 A_n :



范继淹的并立扩展法解决了部分判定问题,比如:

词类序列及层次	实例层次	并立扩展方式
副·动 宾	别说 这样的事儿 不了解 情况	别说别想这样的事 不了解、不调查情况
副 动·宾	先 看电视 很 伤感情 一起 写信	先看电视、听音乐 很 伤感情、伤面子 一起 写信、听音乐
副·动 了	别哭 了	别哭别吵了

	不读 了	不读、不写 了
	太淘气 了	太淘气、太顽皮了
	最清楚 了	最清楚、最明白了
副 动·了	都 读了	都 读了、写了
	又 修了	又 修了、补了

但还有些情况不能判定。如何提出一个更一般的判定原则,需要进一步研究。

11.4 两种递归方式

前面我们讨论过,递归性只体现在扩展替换中,不体现在转换中,由于 X 杠理论本质上是对扩展替换的描写,如果 X 杠中没有递归性,整个句法描写就不可能有递归性。

为了认识一个改写规则或替换模式是否有递归性,需要严格区分递归序列中的递归模式和非递归模式。下面是一个递归序列:

- 两个学生
- 两个一年级的学生
- 两个一年级的学习很好的学生

这个序列是可以无限递归下去的,但能够递归的是带“的”字结构的名词短语,并不是数量名短语。

在改写规则中,递归性的本质体现为直接递归和间接递归。在直接递归中,改写前后都应该有相同的非终端符号:

$A \rightarrow AB$

通过重复使用这条规则就可以生成无限的组合：

$$AB \rightarrow [AB]B \rightarrow [[AB]B]B \dots$$

间接递归体现在不同的改写规则两端出现相同的符号：

$$A \rightarrow BC$$

$$B \rightarrow AD$$

把后一个公式代入前一个公式，立即可以得到直接递归公式：

$$A \rightarrow BC \rightarrow [AD]C$$

递归可以分成核心成分递归和边际成分递归。考虑下面的扩展替换：

[民族的]习俗

→[山区民族的]习俗

→[[高寒山区]民族的]习俗

第一个扩展替换是在“习俗”的修饰语中的“民族”上进行，第二个扩展替换是在修饰语“山区”上进行。如果对修饰语的扩展替换能够持续下去，就形成递归性。这种递归属于附加成分递归或边际递归。下面的扩展替换是不断在短语的核心上进行：

民族的[习俗]

→民族的[独特的习俗]

→民族的(独特的[有保留价值的习俗])

→民族的〈独特的(有保留价值的[有待深入研究的习俗])〉

这种扩展从另一个角度看也可以看成是边际语“民族的”作为整体的并列扩展:

[民族的]习俗

→([民族的][独特的])习俗

→([民族的][独特的][有保留价值的])习俗

→([民族的][独特的][有保留价值的][有待深入研究的])
习俗

尽管如此,边际成分的并列扩展和边际成分中某个成分的扩展仍然是有区别的,边际成分的并列扩展仍然可以理解成对核心的扩展,但边际成分中某个成分的扩展不可以理解成对核心的扩展,所以我们把边际成分的并列扩展仍然看成是对核心的扩展。

以上两种扩展方式重叠在一起就有:

([高寒山区]民族的)〈独特的(有保留价值的[有待深入研究的])〉习俗

一个短语的各个部分都存在扩展替换的可能,无论具体情况多复杂,无非是在修饰成分上扩展和在核心上扩展。由于递归就是无限扩展,所以递归性基本上可以概括成修饰成分的递归和核心的递归。

11.5 X 杠的递归性

根据 Chomsky(1965)的 X 杠理论有:

$$X^2 \rightarrow \text{Spec } X^1$$

$$X^1 \rightarrow X \text{ Comp}$$

可以生成的基本组合是：

$$[X^2 \text{ Spec } [X^1 X \text{ Comp}]]$$

这种 X 杠理论从本质上看只承认短语有两种，一种是 X^2 ，即 $[\text{Spec } X^1]$ ，一种是 X^1 ，即 $[X \text{ Comp}]$ ，我们把这种语杠理论称为 2 阶 2 杠 X 杠理论，2 阶是指短语的种类为 2。下面以名词短语为例来分析 2 阶 2 杠理论在处理递归性方面遇到的问题。考虑英语介词短语修饰名词的情况，按照 2 阶 2 杠 X 阶理论，这时名词短语的改写规则可以是：

$$N^2 \rightarrow \text{Spec } N^1$$

$$N^1 \rightarrow N [\text{Prep } N^2]$$

这时候修饰核心 N 的补足语是一个介词短语。这两个规则反复使用是有递归性的：

$$\text{Spec } N^1$$

$$\rightarrow \text{Spec } N^1 \text{ (使用规则 } N^2 \rightarrow \text{Spec } N^1 \text{)}$$

$$\rightarrow \text{Spec } [N [\text{Prep } N^2]] \text{ (使用规则 } N^1 \rightarrow N [\text{Prep } N^2] \text{)}$$

$$\rightarrow \text{Spec } [N [\text{Prep } [\text{Spec } N^1]]] \text{ (使用规则 } N^2 \rightarrow \text{Spec } N^1 \text{)}$$

.....

可以看出，这种递归性相当于我们前面讨论的边际成分的递归，不是核心成分的递归，因此 2 阶 2 杠的 X 杠理论无法完成下面的生成过程：

the decision

→ the decision on the boat

→ the decision on the boat with my friends

因为这里的生成过程是对核心的不断替换, 实际语类层次是:

Spec ([X Comp] Adjunct)

同样, 2 阶 2 杠 X 杠理论也不能生成下面的汉语片断:

那把雨伞

→ 那把透明的雨伞

→ 那把透明的小雨伞

因为这里的语类关系是:

Spec(Adjunct [Comp X])

由此可见, 2 阶 2 杠的 X 杠理论不能体现补足语和附加语的区别。我们后面会讨论, 由于带附加语的名词短语才可能有核心扩展, 因此 2 阶 2 杠的 X 杠理论只能描写边缘成分的递归性, 不能描写核心成分的递归性。Chomsky(1965) 虽然已经注意到了补足语和附加语的区别, 但 Chomsky 后来的 2 杠的杠理论并没有在 X 杠理论中体现这种区别。

补足语和附加语的区别是明显存在的, 名词短语至少要有 3 个阶, 即带限定成分的名词短语、带附加成分的名词短语和带补足成分的名词短语。要区分这 3 个阶, 一个可能的办法是 X 杠理论中再多一杠, 形成如下三杠的 X 杠模式:

1. $X^3 \rightarrow \text{Spec } X^2$
2. $X^2 \rightarrow X^1 \text{ Adjunct}$
3. $X^1 \rightarrow X \text{ Comp}$

对于具体的动词短语来说,就有:

1. $V^3 \rightarrow \text{Spec } V^2$
2. $V^2 \rightarrow V^1 \text{ Adjunct}$
3. $V^1 \rightarrow V \text{ Comp}$

对于名词短语来说,就有:

1. $N^3 \rightarrow \text{Spec } N^2$
2. $N^2 \rightarrow N^1 \text{ Adjunct}$
3. $N^1 \rightarrow N \text{ Comp}$

根据这里的名词短语规则,下面短语中的补足语和附加语可以得到区分:

$\langle_{N^3} \langle_{\text{Spec}} \text{the} \rangle \langle_{N^2} (\langle_{N^1} [\text{N decision}] [\text{Comp on the boat}]) (\text{Adjunct with my friends}) \rangle \rangle$

上面汉语实例也可以得到解释:

$\langle_{N^3} \{ \text{Spec 那把} \} \{ N^2 (\text{Adjunct 透明的}) (\langle_{N^1} [\text{Comp 小}] [\text{N 雨伞}]) \} \}$

但是,这种依靠添加杠的方法来展开形式化描写尽管能够区分开补足语和附加语,但仍然不能体现核心的递归性,因为附加成分增加,语杠就得增加。比如,如果我们在以上实例上再增加一

个附加语 Adjunct:

the decision [on the boat] [with my friends] [in the class-room]

是否需要把语杠增加到 N^1 ? 类似的, 下面的汉语是否要用 X^5 ?

那把透明的小巧的刚从上海买的小雨伞

从理论上说, 短语是有递归性的, 如果随着修饰成分的增加语杠也需要增加, 这样的短语规则就没有递归性。还有, 以上 3 杠和 n 杠的 X 杠理论尽管能在有限的范围把层次表现出来, 把补足语和附加语区分开, 但由于没有递归性, 并不能表现出补足语 (Comp) 和附加语 (Adjunct) 的根本区别。

自然语言有递归性, 对自然语言进行形式化描写, 必须要有公式来体现递归性。前面我们讨论过, 递归性体现在叠加和替换中, 而短语结构规则是对叠加和替换的形式化描写, 所以短语结构规则要体现出递归性。以上 3 杠的语杠理论由于每个改写规则左右两边都没有相同的符号, 所以并没有核心递归。其递归性只存在于边际成分中, 比如:

1. $N^3 \rightarrow \text{Spec } N^2$
2. $N^2 \rightarrow N^1 \text{ Adjunct}$
3. $N^1 \rightarrow N[\text{prep } N^3]$

即外围成分包含有 N^3 时, 才可能有边际递归。现在的问题是, 怎样才能够在 X 杠理论中恰当的引入递归改写规则, 同时反映边际成分的各种性质。Radford(1988, p177) 给出过名词短语 2 杠的 N 杠规则:

1. $N^2 \rightarrow \text{Spec } N^1$
2. $N^1 \rightarrow N^1 \text{ Adjunct}$
3. $N^1 \rightarrow N \text{ Comp}$

由于该 2 杠理论区分了三种短语,我们可以称之为 3 阶 2 杠理论。这个模式既能够把补足语(Comp)和附加语(Adjunct)区别开,又有递归性。这个模式的含义可以这样解释: N^2 直接统辖(dominate)限定语(Spec)和 N^1 ,或者 N^2 的直接成分是限定语和 N^1 ; N^1 可以直接统辖 N^1 和附加语,也可以直接统辖 N 和补足语,或者说 N^1 的直接成分或者是 N^1 和附加语,或者是 N 和补足语。由于 N^1 可以直接统辖 N^1 和附加语,这一模式就有了递归性。再来看上面提到的 Chomsky 的例子,根据可递归 3 阶 2 杠的 X 杠理论,我们可以作以下分析:

【 $N^2 \langle \text{Spec the} \rangle \langle N^1 (N^1 [N \text{ decision}] [_{\text{Comp}} \text{ on the boat}]) (_{\text{Adjunct}} \text{ with my friends}) \rangle \rangle$ 】

概括各类短语,X 阶理论应该修正为:

- $$\begin{aligned} X^2 &\rightarrow \text{Spec } X^1 \\ X^1 &\rightarrow X^1 \text{ Adjunct} \quad (\text{可选规则}) \\ X^1 &\rightarrow X \text{ Comp} \end{aligned}$$

以上规则就具有核心递归性,可以在 X^1 上不断递归,比如:

spec [$X^1 \text{ Adjunct}$]
spec [[$X^1 \text{ Adjunct}$] Adjunct] (再使用一次 $X^1 \rightarrow X^1 \text{ adjunct}$)

spec [[[X Comp] Adjunct] Adjunct] (使用 $X^1 \rightarrow X$ Comp)

如果这里的 X 是 N, 通过词汇插入就有:

the teacher in the classroom

the teacher with glasses in the classroom

the teacher of biology with glasses in the classroom

作为汉语名词短语的一个实例, 考虑到汉语的修饰语在中心语前面, 我们可以有:

$N^2 \rightarrow \text{Spec } N^1$

$N^1 \rightarrow \text{Adjunct } N^1$

$N^1 \rightarrow \text{Comp } N$

可以生成下面的组合:

这张桌子

这张干净的桌子

这张干净的漂亮的桌子

这张干净的漂亮的玻璃做的桌子

这张干净的漂亮的玻璃做的有古代特点的桌子

这张干净的漂亮的玻璃做的有古代特点的小桌子

X 杠理论到底需要多少杠? Chomsky 根据最大投射认为 X^2 就够了。从我们以上的分析可以看出, 不同的语言, 语杠可能并不一样。对于名词短语来说, 如果加上一个反映递归性的改写规则, X^2 基本上是不够的, 但应该是 3 阶 2 杠而不是 2 阶 2

杠。也有的学者,比如 Jackendoff,认为可以有 X^n ,其中 n 可以无限大,这是从层次的多少来说的。但是,如果没有递归性的改写规则,杠数 n 再多,也不能反映结构的递归性。相反,如果有了递归性改写规则, X^n 无论嵌套多少层,都可以通过改写规则生成出来,所以在改写规则中由于要有递归性规则, X 杠至少要有 1 杠,于是有:

$$X^1 \rightarrow X^1 YP$$

$$X^1 \rightarrow X ZP$$

如果没有第一个规则,则没有递归性,如果没有第 2 个规则,则不可能完成词汇插入。因此 X 杠理论至少要有一杠才能描写自然语言。或许有些自然语言一杠就够了,但从前面汉语和英语的三种修饰语看,由于 NP 必须有 2 杠,如果要概括各种短语的性质, X 杠理论至少要有 2 杠。

11.6 短语结构的阶和语杠理论

描写短语结构最多需要多少杠?从 Chomsky(1965,1970)的研究看,可以分出两杠。但边际成分的种类有好几种,包括补足语(complement)、限定成分(determiner)、附加语(adjunct)等,在两杠中描写并不充分。其实语杠理论和替换的阶是紧密联系的,前面我们分析 X 杠理论时,已经涉及短语结构在替换上的层阶属性。短语结构的阶是人类语言的重要属性, X 杠理论已经涉及阶的问题,但在讨论递归、替换和阶的关系这一根本问题上,还需要深入展开研究。下面我们来进一步讨论这个问题。

前面说对短语的核心进行扩展替换可以生成新的结构,但是,对短语的核心进行扩展替换是有条件的。比如:

小桌子→小[方桌子] (用“方桌子”替换“桌子”, 成功替换)

方桌子→·方[小桌子] (用“小桌子”替换“桌子”, 不成功替换)

干净的桌子→干净的[小桌子] (用“小桌子”替换“桌子”, 成功替换)

小桌子→·小[干净的桌子] (用“干净的桌子”替换“桌子”, 不成功替换)

可见, 同类扩展替换是有条件的。朱德熙(1982)把汉语名词短语分成黏合式名词短语和组合式名词短语, 并认为黏合式短语可以替换组合式短语的中心名词, 但组合式短语不能替换黏合式短语的核心。朱德熙还根据修饰语的种类进一步把组合式偏正结构分成: (1) 定语带“的”的偏正结构(下面记为[XPd]N), (2) 由数量词(或指示代词加量词)做定语的偏正结构(下面记为[QP]N), (3) 表示领属关系的偏正结构(下面记为[PosP]N)。现在我们在此基础上进一步讨论名词短语的替换性质以及阶和杠的一般性质。

名词短语的替换问题从另一个角度看是线性叠加的顺序问题, 即不同类型的修饰语线性位置不同。在上面的实例中, 带“的”的形容词必须在不带“的”的单音节形容词前面。“小”必须在“方”前。朱德熙(1982)关于名词修饰语的顺序也是先从线性顺序开始讨论, 然后统一到替换理论上。后来有学者进一步对名词修饰语的顺序进行过线性解释。为了看出组合的层次, 说明阶和杠的关系, 加强生成的操作性, 我们下面仍然从替换的角度考虑问题。

一般地说, 对于[XP]N 和[YP]N 两个名词短语, 如果:

* $[XP]N \rightarrow [XP][[YP]N]$ (即 $[YP]N$ 不能替换 $[XP]N$ 的核心 N)

$[YP]N \rightarrow [YP][[XP]N]$ (即 $[XP]N$ 能够替换 $[YP]N$ 的核心 N)

我们就说短语 $[YP]N$ 的阶高于短语 $[XP]N$ 的阶。比如：

* $[小]桌子 \rightarrow [小][[两张]桌子]$
 $[两张]桌子 \rightarrow [两张][[小]桌子]$

即“小桌子”中的核心“桌子”不可以被“两张桌子”替换，但“两张桌子”的核心“桌子”可以被“小桌子”替换，我们说“两张桌子”的阶高于“小桌子”的阶。阶的存在是自然语言中的可观察事实，X 杠理论中的杠则是对阶的形式化描述。

从以上实例看，我们很容易分出很多不同的短语阶，但有些阶有根本性的区别。带“的”字修饰语的名词短语 ($[XP_d]N$) 之间 (不包括领属短语)，一个短语的核心可以被另一个短语自由的替换：

木头的桌子
 木头的 { 干净的桌子 }
 木头的 { 干净的 (刷过漆桌子) }
 木头的 { 干净的 (刷过漆的 [古色古香的桌子]) }

替换的顺序也不受限制：

干净的桌子
 干净的 [木头的桌子]
 干净的 [木头的 [刷过漆的桌子]]

干净的[木头的[刷过漆的[古色古香的桌子]]]

如果记忆和语境条件允许,“[XPd]桌子”可以无限次的替换“桌子”。这类短语显然具有递归性。可以概括为:

$$N^1 \rightarrow [XPd] N^1$$

其中 XP 是可以出现在“的”前的短语和单位,但 XPd 不包括领属修饰语。

XPdN 短语属于可无限扩展短语。显然,递归性短语必须是可无限扩展短语。后面会讨论,可无扩展短语不一定是递归性短语,比如[QP]N 和 PosN 都是可无限扩展的,但没有递归性,因为这类短语只能被其他类型的短语替换,但不能用和自己同类型的短语替换:

两只猫

→两只[小猫]

→两只(白色的[小猫])

六只猫

→*六只[两种猫]

他的猫

→*他的[我的猫]

这里的语义解释是,名词只能受一次数量短语或代词短语修饰。

修饰成分不带“的”的名词性短语,即朱德熙(1982)提到的黏合式名词短语,朱德熙认为内部之间是可以替换的,我们注意到这种替换次数是相当有限的:

木头桌子

木头[小桌子]

木头[小[方桌子]]

[木头[小[方桌子]]]

而且内部顺序是相当严格的,以上实例的替换顺序一旦改变,就不成立,或很难接受:

小桌子→*小[木头桌子]

小桌子→*小[老式桌子]

方桌子→*方[小桌子]

方桌子→*方[木头桌子]

这种替换顺序上的限制从另一个角度看就是黏合式短语的修饰语的顺序问题。从替换和递归的角度看,我们说黏合式短语是非递归短语。非递归短语可以替换递归短语的核心:

干净的桌子→干净的[小桌子]

但递归短语不能替换非递归短语的核心:

*小桌子→小[干净的桌子]

这种现象朱德熙(1982)在讨论定语的语序时已经注意到,当时他认为黏合式短语可以替换组合式短语的核心,但没有讨论递归问题。

以上讨论已经涉及了语言的递归性的本质:其核心成分能被不断扩展替换的短语是核心可递归短语,核心成分不能被不断替换的短语是非核心递归短语。于是在短语中首先可以分出

两个大的阶,即核心可递归短语阶和核心不可递归短语阶。从X杠理论形式化描写看,要表现出这两类名词短语的区别,可以建立下面两个规则:

$$N^1 \rightarrow [XPd] N^1 \quad (\text{“XP 的”为带“的”尾的修饰语})$$

$$N^1 \rightarrow [X] N \quad (X \text{ 为单个语符})$$

这两个规则可以把递归性名词短语和非递归性名词短语区别开,也能把组合式名词短语和黏合式名词短语区别开。不过以上的“ $N^1 \rightarrow [X] N$ ”只能用一次就必须进行词汇插入,不能反映出黏合式名词短语可以有限扩展的性质。先观察黏着式名词短语有限扩展的情况:

〈老式{木头(小[方桌子])}〉

为了体现黏合式短语内部能够进行 n 次有限扩展替换的性质,我们把以上规则修正如下:

$$N^{n+1} \rightarrow [XPd] N^{n+1} \quad (\text{“XP 的”为带“的”尾的修饰语})$$

$$N^{n+1} \rightarrow [X] N^{n+1} \quad (X \text{ 为单个语符}, n \geq 0, n \text{ 是有限的})$$

$$N^{n+1} \rightarrow [X] N^n$$

n 的取值表示黏合式名词短语最多可以作 n 次替换,最少可以作 0 次替换。

汉语中黏合式名词短语内部的修饰成分在顺序上比较严格,已经有学者作过研究,找出了一些规则,但要完全形式化,现在找到的规则还不够。

为了后面讨论的方便,我们把“小桌子”($[X] N$)这样的黏合式名词短语称为一阶短语,把“干净的桌子”($[XPd] N^1$)这样

的短语称为二阶短语,于是我们可以总结出这样的性质:短语的核心可以被同阶短语或低阶短语扩展。这一性质蕴含了二阶短语可以无限扩展。

再来看内部带数量短语的名词短语(数量名短语)的情况:

两张桌子

这张桌子

这两张桌子

由于语义上限制,名词只能受一次数量短语修饰,所以数量名短语的核心不能够再被数量名短语替换。下面两个实例不是核心替换关系:

这张桌子

这两张桌子

即并不是用“两张桌子”替换“这张桌子”中的“张桌子”,因为“张桌子”不是符合语法的组合,更不是核心。这里是用“这两张”替换“这张”,不是核心成分的扩展替换,是卫星成分或边际成分的扩展替换。于是我们说数量名短语自身没有递归性。但数量名短语的核心可以被上面提到的二阶短语(递归短语[XP_d]N)替换,这一点和一阶短语(非递归短语)不同:

这张[桌子]

这张(干净的[桌子])

这张(干净的<漂亮的[桌子]>)

这张(干净的<漂亮的{玻璃做的[桌子]}>)

不过我们并不能作出如下形式化描写:

$N^1 \rightarrow [QP] N^1$ (QP: 数量、指量短语)

因为这样会导致数量叠加的组合出现:

$[QP [QP [QP [QP N^1]]]]$

可见,可扩展和可以递归是两个不同的概念,可递归的短语一定是可无限扩展的,如前面的二阶短语“漂亮的桌子”,但可无限扩展的短语不一定是递归的,“数量名”就是这样的短语。递归的本质是:XP 短语的核心能够被 XP 类型的短语替换。

除了上面提到的二阶短语($[XPd]N^1$)可以替换数量名短语的核心,数量名短语($[QP] N^1$)在一定条件下也可以扩展二阶短语的核心:

干净的桌子 $\rightarrow$ 干净的这张桌子

但是,一旦这样替换以后,短语的核心就不能再被二阶短语替换:

干净的这张桌子 $\rightarrow$ * 干净的这张漂亮的桌子

这一性质从线性角度看就是数量短语(QP)通常不能把 APd 跟核心 N 分开:

$[_{APd}$ 干干净净的 $][_{APd}$ 又黑又亮的 $][_{QP}$ 那两张 $]桌子$

$[_{QP}$ 那两张 $][_{APd}$ 干干净净的 $][_{APd}$ 又黑又亮的 $]桌子$

* $[_{APd}$ 干干净净的 $][_{QP}$ 那两张 $][_{APd}$ 又黑又亮的 $]桌子$

* $[_{APd}$ 又黑又亮的 $][_{QP}$ 那两张 $][_{APd}$ 干干净净的 $]桌子$

QP 可以出在 APd、VPd 之间,但不是很合适:

[_{QP} 那两张][_{VPd} 手工做的][_{APd} 又黑又亮的]桌子
 [_{VPd} 手工做的][_{APd} 又黑又亮的][_{QP} 那两张]桌子
 ? [_{VPd} 手工做的][_{QP} 那两张][_{APd} 又黑又亮的]桌子
 ? [_{APd} 又黑又亮的][_{QP} 那两张][_{VPd} 手工做的]桌子

QP 可以把 VPd 序列分开,但也不是很好的实例:

[_{QP} 那两张][_{VPd} 手工做的][_{VPd} 刷过漆的]桌子
 [_{VPd} 手工做的][_{VPd} 刷过漆的][_{QP} 那两张]桌子
 ? [_{VPd} 手工做的][_{QP} 那两张][_{VPd} 刷过漆的]桌子
 ? [_{VPd} 刷过漆的][_{QP} 那两张][_{VPd} 手工做的]桌子

从总体上看,尽管 QP 可以处在 XPd 序列的前面、中间或后面,但毕竟有一定的语类限制。这说明[QP]N 名词短语和[XPd]N 名词短语在替换性质上有区别,两者的替换存在这样一种顺序,当几个[XPd]N 之间的核心替换程序完成以后,一旦使用[QP]N 名词短语替换[XPd]N 名词短语的核心,就不再能使用[XPd]N 来替换短语的核心,或者说替换起来有一定的条件。从线性叠加的角度看,后面带“的”字的附加语叠加到名词前以后,如果再叠加数量短语,则前面继续叠加“的”字短语会受到一定限制。

如果把[QP]N 勉强看成和[XPd]N 同阶的名词短语,并依照上面的语杠理论来描写:

$$N^{n+1} \rightarrow [YP] N^n \quad (YP \text{ 包括 } QP \text{ 和 } XPd)$$

$$N^{n+1} \rightarrow [X] N^n \quad (X \text{ 为单个语符}, n \geq 0, n \text{ 是有限的})$$

不仅可能出现前面提到的数量叠加的非法组合,而且可能生成这样的组合序列:

- 干干净净的那两张又黑又亮的桌子
- 又黑又亮的那两张干干净净的桌子

因此,从替换条件看,应该把[QP]N看成和[XPd]N不同阶的名词短语。

当然也不能把数量名短语[QP]N看成是和XN相同的1阶名词短语,因为[QP]N的核心是可以无限扩展的,所以下面规则会存在问题:

$$\begin{aligned} N^1 &\rightarrow [XPd] N^1 \\ N^1 &\rightarrow [QP] N \\ N^1 &\rightarrow [X] N \end{aligned}$$

因为一旦使用了第2个规则,第3个规则就都不能使用了,这 and 实际存在的现象不一致:

崭新的[书]→崭新的[一本书]→崭新的[一本(学术书)]

也就是说,以上改写规则不能反映出[QP]N名词短语和XN名词短语在阶上的区别,归根结底是因为[QP]N的中心是可以用[XPd]N¹和[X]N短语无限扩展的:

这张桌子
 这张{旧的桌子}
 这张{旧的<小[桌子]>}
 这张{旧的<小[方桌子]>}

这张{旧的<小(木头[方桌子])>>}

这要求有这样的组合:

{QP<XPd(X[N])>>}

但前面的形式化描写不能生成这样的组合。因此,从严格意义上可以说[QP]N 名词短语和[XPd]N 名词短语不是一个替换阶上的短语。

为了要在语杠描写中把这两种阶区别开,同时反映出[QP]N 名词短语是可无限扩展不递归的,[XPd]N 名词短语是可递归的,我们可以尝试考虑下面的两种描写方式:

$N^2 \rightarrow [XPd] N^2$ (XP 可以是动词短语、名词短语、形容词短语,但不能是介词短语)

$N^2 \rightarrow [QP] N^1$

$N^1 \rightarrow [X] N$ (X 可以是能够直接修饰名词的语符)

$N^2 \rightarrow [QP] N^1$

$N^1 \rightarrow [XPd] N^1$

$N^1 \rightarrow [X] N$

这两种描写中,都只有[XPd] N^i (i 为 1 和 2) 可以是递归的。这两组规则可以分别生成下面两个序列:

<XPd(QP[XN])>

<QP(XPd[XN])>

那么[QP]N 名词短语和[XPd]N 名词短语到底哪个阶更高?

汉语中我们还没有找到充分的证据。不过下面左列的实例似乎更常用,更能够说明数量名短语的阶更高,因为 XPd[[QP]N]是相当受限制的:

QP[[XPd]N]	XPd[[QP]N](使用有条件)
两张干净的桌子	干净的两张桌子
两只调皮的猫	调皮的两只猫
两只黑色的眼睛	黑色的两只眼睛

另一方面,从类型学的角度看,很多语言似乎都是指称短语(和[QP]N 相当)比描述短语(和[XPd] N 短语相当)的阶更高,即描述短语可以替换指称短语的核心,但指称短语通常不能替换描述短语的核心,比如英语就属于这类情况:

the [teacher] → the [English teacher]
English[teacher] → * English [the teacher]

因此,如果要从理论上进行统一,可以考虑把数量名短语的阶看成是比描述短语的阶要高。最终选择这样的描写规则:

$N^2 \rightarrow [QP] N^1$
 $N^1 \rightarrow [XPd] N^1$
 $N^1 \rightarrow [X] N$

这样的规则能够生成“这张干净的桌子”,也可以避免“* 干净的这张漂亮的桌子”出现,但不能生成“干净的这张桌子”。我们可以把“干净的这张桌子”这种形式看成是转换的结果,转换的条件是语义和语用的细微差别。关于这种差别现在还没有得到充分描写,能够观察到一些:

我选了这张干净的桌子	? 我选了干净的这张桌子
这张干净的桌子缺一条腿	? 干净的这张桌子缺一条腿
这张干净的桌子旁边有一把椅子	? 干净的这张桌子旁边有一把椅子

右列的分布是不太容易接受的,或者说需要有特殊的条件。

现在来看领属定语构成的名词短语:

我的送给张三的那本描写农村生活的很严肃的小说

由于语义条件的限制,一个名词短语通常只有一个领属成分,所以领属短语的核心不能再被领属短语扩展,领属短语因此也没有递归性。

普通名词、专名、人称代词作定语有一定的区别,比较:

		领属短语 的性质	核心的 性质
图书馆的两本书	两本图书馆的书	普通名词+的	普通名词
桌子的两条腿	两条桌子的腿	普通名词+的	普通名词
鲁迅的两本书	两本鲁迅的书	专名+的	普通名词
张三的两本书	两本张三的书	专名+的	普通名词
张三的两个哥哥	? 两个张三的哥哥	专名+的	亲属词
我的两个哥哥	? 两个我的哥哥	人称代词+的	亲属词
你的两个哥哥	? 两个你的哥哥	人称代词+的	亲属词
我两个哥哥	两个我哥哥	人称代词	亲属词
北京大学两位老师	两位北京大学老师	专名	普通名词
北京大学的两位老师	两位北京大学的老师	专名+的	普通名词

“张三两个弟弟”不是“张三弟弟”扩展来的,因为“张三弟弟”不

成立。更一般地说,汉语中专名不带“的”不直接修饰亲属词。“张三两个弟弟”可以看成是“张三的两个弟弟”省略了“的”。“张三两个弟弟”还有一种主谓关系,不属于这里讨论的情况。

领属短语的核心可以被其他名词短语扩展替换:

我[弟弟]→我[两个弟弟]

我的[弟弟]→我的[两个弟弟]

我[弟弟]→我[喜欢足球的弟弟]

我的[弟弟]→我的[喜欢足球的弟弟]

张三的[弟弟]→张三的[两个弟弟]

张三的[弟弟]→张三的[喜欢足球的弟弟]

如果领属成分是名词带“的”,这时的修饰成分具有“的”字修饰语的性质,因此这时的领属短语可以替换数量名短语和 XPdN 短语的核心,扩展的顺序也比较自由:

两本书→两本[鲁迅的书]→两本[鲁迅的[再版的书]]

再版的书→再版的[鲁迅的书]→再版的[鲁迅的[两本书]]

很畅销的书→很畅销的[鲁迅的书]

但如果领属短语的修饰语是代词或代词加“的”,一般不能替换其他短语的核心:

两个弟弟→*两个[我弟弟]

喜欢足球的弟弟→*喜欢足球的[我弟弟]

两个弟弟→*两个[我的弟弟]

喜欢足球的弟弟→*喜欢足球的[我的弟弟]

可见以代词作领属成分的名词短语处在最高阶,而以名词作领

属成分的名词短语可以扩展数量名短语和 XPdN 短语的核心,也可以被数量名短语和情状短语扩展,和代词或代词加“的”作修饰语的名词短语在替换性质上差别比较大。考虑到上面对数量短语的分析,在汉语的名词短语中,至少应该有四阶名词短语:

1. 分类性短语,修饰语不带“的”的名词、形容词、区别词充当,内部扩展替换是有限的
2. 描述性短语,除了人称代词以外修饰语加“的”的名词短语,是递归的
3. 指称性短语,包括数量短语,其核心可被低阶短语扩展,但指称性短语本身不递归
4. 代词领属短语,其核心可被低阶短语扩展,但代词领属短语不递归

根据上面分析的扩展替换规律,我们可以进一步总结出这样的性质:短语的核心可以被同阶短语或低阶短语扩展。具体地说,在汉语名词短语中,四阶短语可以被三阶短语、二阶短语和一阶短语扩展,三阶短语可以被二阶短语和一阶短语扩展,二阶短语可以被二阶短语和一阶短语扩展,一阶短语可以被一阶短语扩展。在名词短语的四个阶中,二阶短语之间可以无限扩展,具有递归性。

要体现领属名词短语的阶,还需要有 3 杠:

$$N^3 \rightarrow [\text{Pro}] N^2 \quad (\text{Pro: 代词领属短语})$$

$$N^2 \rightarrow [\text{QP}] N^1$$

$$N^1 \rightarrow [\text{XPd}] N^1 \quad (\text{XPd: “的”字短语})$$

$$N^1 \rightarrow [\text{Atr}] N \quad (\text{Atr: 属性修饰语})$$

考虑到下面英语的实例,英语名词短语的杠似乎也应该是 3 杠,因为下面四个阶是存在的:

〈领属短语 his {数量名短语 three (描述短语 [分类短语 students of physics] in classroom)}〉

当然,我们可以不改变生成语法理论 2 杠的统一理论模式,但必须再附加其他规则来解释 3 阶名词短语(指称性短语)和 4 阶名词短语(代词领属短语)的区别,可见名词短语的阶是语杠理论的基础。

以上我们把[XPd]N 名词短语看成递归性短语,把[QP]N 和[Pro]N 看成非递归短语,这已经涉及语义和逻辑问题,因为名词短语从逻辑上说不可以有两个或两个以上的数量修饰,不可以有两个或两个以上的领属修饰。下面的多个领属成分的出现应该是并列替换的结果而不是递归替换的结果:

这本书是[_{pro}我的、你的和汪锋的]研究成果。

X 杠理论到底需要几杠,还需要对各种语类短语的阶展开研究。杠越多,对规则的表述越细。比如,在黏合式短语中,我们就可以分出很多次杠:

$$N^1 \rightarrow [\text{双音节词}] N^3$$

$$N^3 \rightarrow [\text{单音节尺寸词}] N^2$$

$$N^2 \rightarrow [\text{单音节颜色词}] N^1$$

$$N^1 \rightarrow [\text{单音节形状词}] N$$

以上规则可以生成:

【_{N⁴} 玻璃〈_{N³} 小{_{N²} 黑(_{N¹} 圆[_N 桌])}>>】

【_{N⁴} 木头〈_{N³} 小{_{N²} 黑(_{N¹} 圆[_N 桌])}>>】

这里的杠还可以增加,不过是有限的。如果把这里的杠和前面描写的组合式名词短语统一起来,杠也会增加:

$N^6 \rightarrow [\text{Pos}] N^5$ (Pos: 领属短语)
 $N^5 \rightarrow [\text{QP}] N^4$
 $N^4 \rightarrow [\text{XPd}] N^4$ (XPd: “的”字短语)
 $N^4 \rightarrow [\text{双音节词}] N^3$
 $N^3 \rightarrow [\text{单音节尺寸词}] N^2$
 $N^2 \rightarrow [\text{单音节颜色词}] N^1$
 $N^1 \rightarrow [\text{单音节形状词}] N$

以上规则可以生成:

|_{N⁶} 汪锋的〈_{N⁵} 两张/_{N⁴} 干净的【_{N⁴} 玻璃〈_{N³} 小{_{N²} 黑(_{N¹} 圆[_N 桌])}>>】/〉|

可以看出,杠的多少取决于描写的宽严。如果语杠描写中采取宽的态度,语法的其他部分就要增加内容,如果语杠描写中采取严的态度,语法的其他部分就会减少内容。不过杠的层次不应该是无限的。Jackendoff(1977)认为语杠理论可以统一描写成以下规则:

$X^n \rightarrow \dots X^{n+1} \dots$ (其中 $n \geq 1$, 并且可以无限大)

这种规则不仅没有递归性,而且,如果有多少次替换就需要多少杠,杠理论就失去了从有限规则到无限组合实例的描写能力。

一般地说,语类规则都应该具有以下形式:

$$X^n \rightarrow \dots X^m \dots (n \geq m, n \geq 1)$$

当 n 等于 m 时,就是一个递归规则。

无论语杠的描写是宽还是严,替换的阶的存在是可观察事实。不同的宽严描写中,带 XPdN 名词短语语杠处理可以不一样,但只有带 XPdN 名词短语有递归性,这种性质是不会改变的。

陆丙甫(1988)在朱德熙(1982)关于组合和黏合的区分的基础上,进一步把黏合式多项定语的语序总结为由三方面的原因限制:(1)越是能使中心语外延缩小的定语越是靠后;(2)越是表示事物稳定性状的黏合定语越靠后;(3)音节少的在后。马庆株(1995)总结的规律是:大小新旧类 > 颜色 > 形状。张伯江、方梅(1996)根据语义类来划分形容词的意义层级,得出的排序规律是“德性、价值、真假、新旧、属性、度量、色彩”。张敏(1998)从认知的角度出发,分析了决定定语顺序的“距离象似性原则”,认为定语和中心语之间的距离取决于它们所表达的概念距离。方希(1999)从语义、语音两个方面给出了几条规则。以上研究大体上是一致的,但也有一些出入,有些属于部分组合正在凝固化,有些属于条件控制的宽严不一样。下面我们把条件控制更严格一些,来看规律所在。先考察单音节语符组合的情况:

小方木凳

小圆木桌

小黑木桌

小黑方桌

小圆黑木桌

小新木伞

旧圆木桌

基本上可以说,在单音节的条件下,黏合式名词短语的修饰语顺序是:

[大小][新旧][颜色][形状][质料]N

由于近距离体现了更常见的分类,所以以上序列在特定的分类活动中会有变化:

黑长袍

圆黑帽

黑短发

短黑发

如果以上某个修饰语被替换成语义相当的双音节语符,这个双音节语符需要转换到最前面:

小圆[木]桌→[木头]小圆桌

小[圆]木桌→[椭圆]小木桌

小方[黑]桌→[黑色]小方桌

小[旧]木桌→[旧式]小木桌

方希(1999)已经从线性排列的角度观察到了在语义等级相同的条件下,双音节或多音节修饰语在单音节修饰语前面的现象,如果语义等级不同,则节律不起作用。我们也可以这样来概括这里规则,语义类相同时双音节或多音节可以转换,语义类不相同同时双音节或多音节不能转换。我们上面的分析说明在不同语义等级之间双音节引起的转换也是可以的。我们的进一步分

析发现,不能转换的往往已经形成固定搭配,即其核心通常不能进行扩展替换。比如:

热洗澡水/小西北风/大黄沙碗/破眼镜盒子/小酒醉螃蟹/
小柿饼帽子/小油漆凳儿/老别扭鬼/破墨水瓶/扁倭瓜脸/
大豆绿蛾/新缎子鞋/破缎子鞋/破近视镜/小独立国/小木头人/
小甜水井/小土地祠/小香烟店/小银钩虾/圆葫芦头/
旧衣裳架子/小广播电台/小山羊胡子/新办公大楼/破闷葫芦罐/
新研究生大楼/旧办公桌子/老石油工人/粘豆沙饼/
男喜剧演员/女纺织工人/小骆驼桥儿/大绿豆蝇/小杂货铺/
新小脚鞋儿/高底洋皮鞋/大黄油鸡卵

像“破墨水瓶”中的“墨水瓶”,其核心通常不能进行扩展替换,所以“墨水”是不能转换到前面,不能有“墨水破瓶”。而像“小木头人”,由于“木头人”的核心还可以进行扩展替换,可以说“木头矮人”、“木头土人”等,所以“小木头人”实际上还可以转换成“木头小人”。

11.7 各种语类的阶和杠

在生成语法的文献中,英语中各种语类都可以统一到 X 杠 2 的理论上来。下面我们把这些语类内部的成分性质都表示出来:

- a. $\langle N^2 (Spec\ a) (N^1 [N\ teacher] [Comp\ of\ Chinese]) \rangle$
- b. $\langle V^2 (Spec\ is) (V^1 [V\ talking] [Comp\ about\ him]) \rangle$
- c. $\langle A^2 (Spec\ very) (A^1 [A\ good] [Comp\ at\ Chinese]) \rangle$
- d. $\langle Adv^2 (Spec\ quite) (Adv^1 [Adv\ differently] [Comp\ from\ me]) \rangle$
- e. $\langle P^2 (Spec\ right) (P^1 [P\ on] [Comp\ of\ Chinese]) \rangle$

在汉语中,用2杠来描写最大投射意义上的名词短语、动词短语和介词短语基本上是可行的:

$\langle N^2 (Spec \text{ 一位}) (N^1 [Comp \text{ 汉语}] [N \text{ 老师}]) \rangle$

$\langle V^2 (Spec \text{ 认真}) (V^1 [V \text{ 教}] [Comp \text{ 汉语}]) \rangle$

$\langle Adv^2 (Spec \text{ 完全}) (Adv^1 [Adv \text{ 关于}] [Comp \text{ 哲学}]) \rangle$ (如:[完全关于哲学]的报告)

对于形容词来说,有些是可行的:

$[A^2 [Spec \text{ 很}] [A^1 [A \text{ 擅长}] [Comp \text{ 汉语}]]]$

有些会出现问题。比如:

$[Comp \text{ 对学习}] [Spec \text{ 很}] [A \text{ 认真}]$

如果把“很”看成是限定语(Spec),这时限定语就离核心“认真”更近,和“认真”互为直接成分。补足语(Comp)“对学习”离核心“认真”比较远,和“很认真”互为直接成分,这和“很擅长汉语”的内在结构有差异,难以在X杠理论上统一起来。X杠理论的基本要求是,补足语和核心X组合成 X^1 ,限定语和 X^1 组合成 X^2 。

汉语中严格意义上的副词一般是指只能作状语的词,比如“很、都、也、就、刚刚、忽然、简直、索性”,这些词并不带卫星成分,因此汉语中可能不存在以这些副词为核心的副词短语,也不存在用X杠2理论来描写的问题。英语中带卫星成分的副词,比如上面的例子 differently,在汉语中都是由形容词充当的,只不过书写中都是用“地”而不是“的”来引入,口语中和形容词短语的读音没有区别:

他_[Adv]跟大家很不一样地]做了一个动作

所以在汉语中这类副词短语性质和形容词短语的性质在本质上是相似的,所遇到的问题也就是形容词短语遇到的问题。

汉语中严格意义上的介词是不能做动词的介词,包括“把、被、连(连水也不喝)、对于、关于、从”,介词短语在汉语中带限定语(Spec)的比较少。

由于不同的语言之间语类的具体内容不同,X 杠理论在不同语言中的统一性问题还需要作深入研究。

前面通过扩展替换讨论过短语的阶,可以看出不是所有的语类短语都有递归性。前面主要讨论的是名词短语的递归性。动词短语也有递归性,动词短语的递归性体现在带有“地”修饰语的动词短语中:

仔细地分析

→仔细地[认真地分析]

→仔细地[认真地[全面地分析]]

→仔细地[认真地[全面地[一个不漏地分析]]]

→仔细地[认真地[全面地[一个不漏地[逐字逐句地分析]]]]

形容词短语似乎没有递归性,因为形容词短语的核心只能进行相当有限的扩展:

高兴→很高兴

高兴→高兴得跳起来

高兴→高兴死了

认真→对工作认真→对工作很认真→对工作认真得难以让

人相信

下面的扩展不是形容词短语核心的扩展,而是补足语内部成分 NP 的扩展,因此不能用来证明形容词短语有递归性。

激动得连〈话〉都说不出来

激动得连〈一句(话)〉都说不出来

激动得连〈一句(感谢的[话])〉都说不出来

如果认为介词短语的核心是介词,由于介词短语的核心是不能扩展的,所以介词短语也没有递归性,其他语言的情况也应该是这样的。下面的扩展是介词所支配的名词短语的扩展,不是介词短语的扩展,或者说是介词的补足语中的 NP 的扩展,而不是介词短语核心的扩展:

从中学

从【北京的中学】

从【北京的〈一所中学〉】

从【北京的〈一所(很好的中学)〉】

从【北京的〈一所(很好的[职业中学])〉】

11.8 X 杠理论的核心问题

在 Bloomfield 的向心结构理论中,核心(head)是可以单说的词和词组,也就是通常意义上的实词。X 杠理论确定这一类核心基本上没有问题。X 杠理论的核心(head)也可以是不能单说的成分,比如虚词、功能词甚至形态,这方面的判定标准会有问题,比如一般情况下限定词(Spec)是卫星成分,但也有人认为是核心,比如把 the milk 中的 the 看成核心。

把什么看成短语的核心,取决于我们要解决什么问题。Bloomfield 区分向心结构和离心结构的概念,是解决成分的分布和短语结构的分布的关系。他认为向心结构的分布特点至少和某个成分的分布相同,而离心结构的分布和成分的分布不同。核心正是在结构分布和成分分布的异同中确定的。这种办法尽管能够在直接成分的分布基础上预测向心短语结构的分布,并确定结构的核心,但并不能预测离心短语结构的分布,因此离心结构也无所谓核心。而且更严格地说,在很多情况下,向心结构的分布和直接成分的分布并不完全相同,比如:

写

写诗

休息

“写诗”的分布并不完全和“写”相同,而是和“休息”相同。另一方面,很多人把主谓关系理解成离心结构,但这样的离心结构往往和它的直接成分的分布有很多相似的地方:

小邵[要来]

由此可见,同与不同是有宽严之分的,用同与不同并不能充分预测结构的分布问题,向心结构和离心结构的界限并不清楚,通过分布的异同来确定结构的核心也会遇到困难。

核心概念并不是不重要,因为在典型的向心结构中,组合的选择条件都跟核心一致。比较:

喝水 喝[矿泉水] *吃[矿泉水]

抽烟 抽[云南烟] *吃[云南烟]

吃菜 吃[四川菜] *喝[四川菜]

我们认为,核心应该以组合的共现条件为标准,这样有利于说明组合关系。考虑下面的实例:

一匹黑马

一只黑猫

选择“马”时,量词要选择“匹”,选择“猫”时,量词要选词“只”,即“马”和量词“匹”共现,“猫”和量词“只”共现,概括地说,当名词和数量短语组合时,名词是和量词有选择关系,因此我们说,数量短语的核心是量词。

介词短语和动词组合时,介词和动词有选择关系,因此介词应该是核心:

从山上下来/*朝山上下来

把介词等作为核心有利于说明是介词造成了介词短语的独立功能,不是名词。从平行关系上看:

从北京/在北京/去北京

都有一种支配关系。另一方面,介词大多数和动词也有历史上的延续关系。

在生成语法中,规定介词也是核心,从选择关系看也有道理的,比较:

stand on the ground/* stand in the ground

再比较下面的英语句子:

He [is coming]

They [are coming]

由于括号中谓语部分的选择限制是落在 is 和 are 上的,所以 is 和 are 是核心,Chomsky(1981)把 is 和 are 看成核心,从选择的标准看是有道理的。

根据扩展替换,可以确定一个多次扩展替换短语的核心:

这张[桌子]

这张[木头[桌子]]

这张[木头[圆[桌子]]]

由于“桌子”是“木头桌子”的核心,所以当“木头桌子”扩展替换“桌子”时,可以看成是“这张木头桌子”的核心。同样道理,“桌子”也是“这张木头圆桌子”的核心。这个实例的特点是核心不断被同一类型的新的结构替换,即不断被偏正结构替换。考虑下面的替换:

小邵[来]

小邵[来[北大]]

由于“来”是“来北大”的核心,所以当“来北大”扩展替换“来”时,“来北大”仍然可以看成是“小邵来北大”的核心。这和传统的中心词分析法有一致的地方。

由于 X 杠理论和最大投射是联系在一起的,XP 的核心也就是最大投射的核心,因此 XP 是在深层结构上讨论成分之间的关系,从这一点可以看出 XP 理论和中心词分析和直接成分分析有一些根本的区别。考虑实例:

一块烤牛排

a roasted beefsteak

在英语中, a roasted beefsteak 由于是周遍性的规则变化, 其生成过程可以描写成

$$[_{VP}[_V \text{ roast}][_N \text{ beefsteak}]] \rightarrow [_{NP} \text{ a roasted beefsteak}]$$

其一般模式是:

$$[_{VP} V + NP] \rightarrow [_{NP} a + Ved + NP]$$

从论元结构的投射层面看, 动词 V 仍然是核心, 和投射原则可以统一起来。汉语中由于动词作定语的周遍条件还没有完全找到, “一块烤牛排”需要处理成:

$$\{[_{NP}(\text{一块})([_V \text{ 烤}][_N \text{ 牛排}])]\}$$

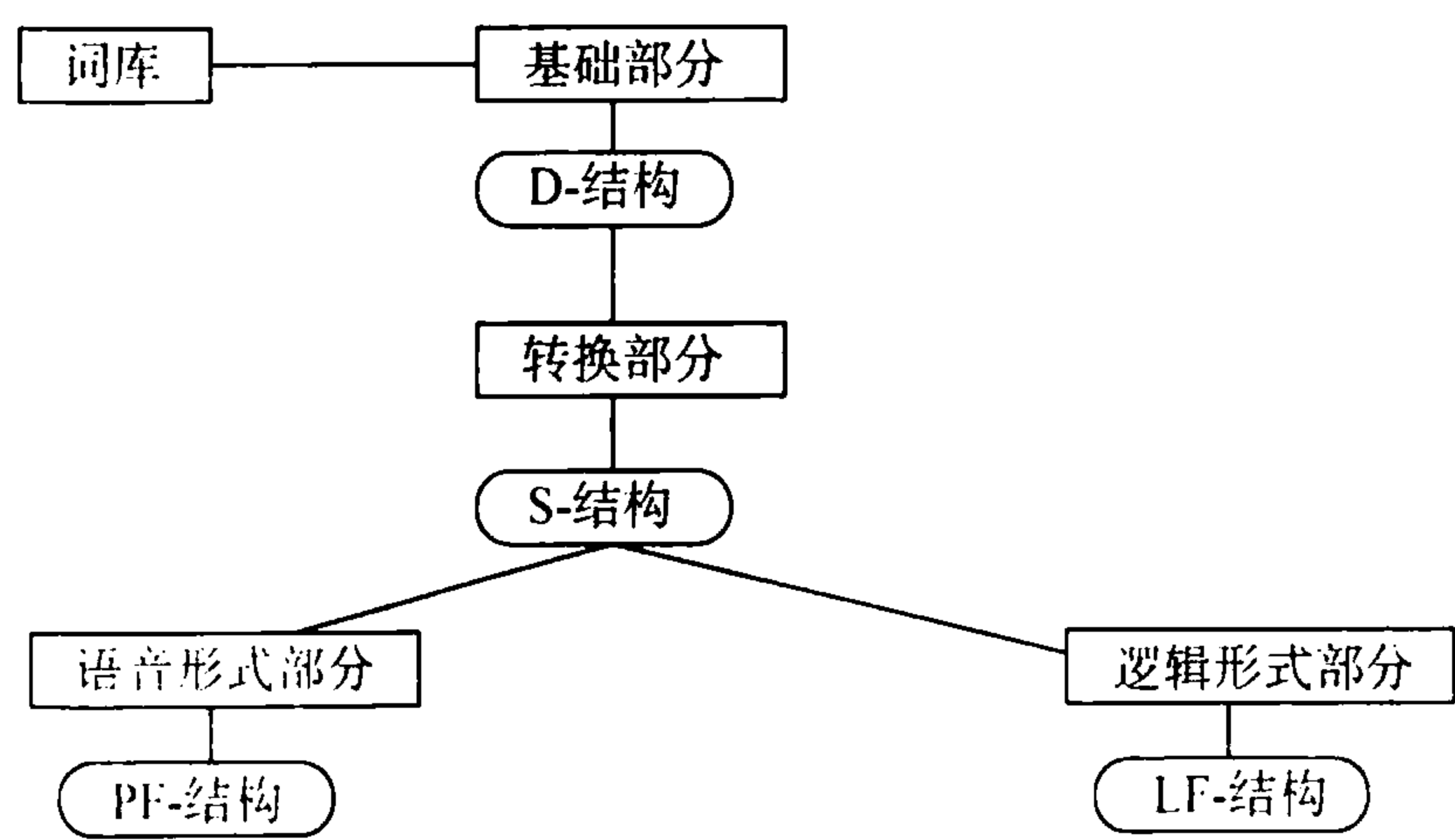
这种现象如何统一起来, 还需要研究。

思考问题

1. 语杠理论存在的问题是什么?
2. 汉语中哪些结构具有递归性?

12. 句法结构关系的初始性

回到本书开头提到的一个重要问题：线性结构是否可以通过单位的分布得到充分的描写？我们说 Harris(1946)从语素到话语的思想就是这个思路，他通过语素的类来解释短语的类。我们前面讨论过，如果把规则单位和不规则单位区别开，把单位建立在语符上，我们基本上可以从形式上判定短语的类并确定语类序列的直接成分。事实上 Wells 就是在 Harris 的基础上确定直接成分的。Harris 当时没有讨论结构关系，Chomsky(1965)的标准理论解释说结构关系或句法结构成分可以通过语类规则推导出来，但 Chomsky 继标准理论以后，在很多研究中都在使用和句法结构关系相关的概念。Chomsky(1981)正式提出支配 (government) 的概念来全面解释各种句法、语义和转换问题。Chomsky(1981) 的 GB (支配约束理论, government-binding theory) 的系统框架是：



D-结构相当于深层结构，S-结构相当于表层结构，但有更严格的

含义。与标准理论相比,GB 的语义解释是在 S-结构部分。由于从词库到 D-结构再到 S-结构要遵循投射原则,即语义结构关系或论元关系保持不变,所以在 S-结构中解释语义得到逻辑形式部分是可行的。

原则子系统有人简称为原则系统。包括以下内容:

- (i) 界限理论(bounding theory)
- (ii) 支配理论(government theory)
- (iii) 题元理论(θ -theory)
- (iv) 约束理论(binding theory)
- (v) 格理论(Case-theory)
- (vi) 控制理论(control theory)

GB 理论的原则部分由于高度抽象和形式化,因此不容易被理解。实际上这个理论模型和生成语法其他理论模型相比,正好证明了句法结构关系的重要性。支配概念作为 GB 的核心,讨论的根本问题是句法结构关系中的支配关系,界限理论主要通过支配概念来定义成分移动的界限,题元理论主要讨论支配关系中论元(argument)角色的指派规则,约束理论主要讨论支配域中论元的同指关系,格理论讨论名词性成分在表层中如何通过支配关系赋格,控制理论讨论支配域中一个动词如何控制另一个不定式动词的潜在主语。

下面我们主要结合约束理论来分析支配(government)的实质,论证句法结构关系的初始性,最终说明线性结构必须要有句法关系的概念。

12.1 支配和约束

支配的概念是支配约束理论的核心。由于成分的移位会受

到结构的线性长度的限制,不允许其中成分移出的句法结构就形成了孤岛,也叫孤岛条件或孤岛限制(island constrain)。孤岛条件主要包括:复合名词词组限制、并列结构限制、主语从句限制和左分支条件等(Ross, 1967)。汉语的移位现象已经得到部分研究,比如话题移位的研究,但基本原则是什么还不清楚。不过从各家使用的材料看,都涉及了句法结构关系。限制移位的长度怎么确定,支配域的研究其实本质上就在考虑这个问题。下面我们要论证,支配域是要以结构关系为条件的。也就是说,结构关系这一功能概念是初始的。

Reinhart(1976)讨论了照应成分(anaphors)受 c-command 的限制的条件,涉及指称的约束和自由问题。约束和自由本来是谓词演算中和量词有关的概念,但自然语言约束与自由的表现方式要复杂得多。支配与约束理论中所说的约束是指某个成分和先于它的某个名词短语指称相同。我们认为研究这个问题具有相当重要的意义,因为尽管前面已经给出了很多关于语义格的分析理论,但面对一个表层句子,一个成分和哪个语义格同指,是一个需要判定的重要问题。Chomsky(1981, 3.1)详细讨论了约束问题,给出了大量实例分析。在英语中,有下面的句子:

John_i saw him_j

John_i saw himself_i

John_i saw John_j

以上有相同右下标的词表示同指(指称相同),右下标不同表示不同指。him 不是指 John, himself 是指 John, 最后一句的两个 John 不是指同一个人。以上差别的一般规律是什么? Chomsky(1981, p188)通过支配范畴给出了约束理论(binding theory)来解释这里的规律。约束理论的核心是约束三原则,都跟支配概念有关系:

(12) Binding Theory

- (A) An anaphor is bound in its governing category
- (B) A pronominal is free in its governing category
- (C) An R-expression is free

翻译成汉语就是：

- (1) 约束原则 A(Binding Principle A)：照应词在支配范畴内受约束(bound)；
- (2) 约束原则 B(Binding Principle B)：代名词在支配范畴内是自由的(free)；
- (3) 约束原则 C(Binding Principle C)：指称词总是自由的

Chomsky(1981,p188)对 governing category(支配范畴)的定义是建立在支配基础上的：

- (11) α is the governing category for β if and only α is the minimal category containing β and a governor of β , where $\alpha = \text{NP or S}$

这里所谓支配范畴,是指最小的 S 或 NP,在这个 S 或 NP 中, α 是 β 的支配范畴的条件是包含并支配 β 的最小范畴。

以上三原则把上面出现的词分成三种：

(1) 照应词(anaphor),包括反身代词 himself, herself 等,也包括相互代词 each other 等；

(2) 代词(pronominal),如 him, her, us, me, them, you 等；

(3) 指称成分(R-expression), 即人名和定指的名词性成分, 如: John, the man

根据约束三原则, 可以比较合理解释以上指称的同一性问题。观察下面实例:

- (1) [_s Tom_i knows himself_i] (原则一, 照应词 himself 在支配范畴(s)内受它的先行语 Tom 约束, 两者同下标。)
- (2) [_s Tom_i knows him_j] (原则二, 代名词 him 在支配范畴(s)内不受约束, Tom 和 him 指称不同, 两者不同下标。)
- (3) [_s Tom_i knows Tom_j] (原则三, 指称词成分 is 自由的, Tom 和 Tom 不同标。)
- (4) [_{s1} Tom_i says [_s Bill_j knows himself_j]] (原则一, himself 在支配范畴内受约束, 必须和 Bill 同标, 而和 S1 中的 Tom 不同标, 因为 Tom 和 Bill 不同标。)
- (5) [_{s1} Tom_i says [_s Bill_j knows him_{i,k}]] (原则二, him 在支配范畴(S)内自由, 所以和 Bill 不同标。但在支配范畴外可以自由, 也可以不自由, 所以可以和 Tom 同标, 也可以不同标。)
- (6) [_{s1} Tom_i says [_s Bill_j knows Tom_k]] (原则三, Tom_k 总是自由的, 所以和 S 中的 Bill 以及 S₁ 中的 Tom 的指称均不同, 不同标。)

除了 S, NP 短语中的照应关系也能由上述三原则来解释:

- (7) [_{NP} John's_i picture of himself_i] (约翰给自己拍的照片)
- (8) [_{NP} John's_i picture of him_j] (约翰给他拍的照片)

(9) [_{NP} John's_i picture of John_j] (约翰给约翰拍的照片)

以上的例句比较好地说明了 Chomsky 关于约束理论三原则的合理性。但是约束理论也存在一些问题,其中主要是关于支配域或支配范畴的定义问题。在后来的工作中,生成语法对支配的概念有各种修正,现在还不能说已经给出了满意的结果,但有一点比较明确,各种修正都是在一个结构的核心与附加的关系范围内讨论支配,这说明关于指称的约束这样一些必须要解释清楚的句法概念,都是以结构关系为条件的。与此类似,Chomsky(1981)的边界理论(bounding theory),也是要以支配为条件的,进一步也是必须以结构关系为前提的,此处不展开。

12.2 支配概念的扩展和长距离约束的传递条件

汉语中的约束现象是否也需要支配概念,从而也需要结构关系的概念?在上述 Chomsky 的约束三原则中,B 原则(代词原则)基本适用于汉语:

张三_i 认识他_j

张三_i 说[李四_j 认识他_{i,k}]

在约束三原则中,C 原则(指称词原则)也基本适用于汉语:

张三_i 认识张三_j

张三_i 说[李四_j 认识张三_k]

在约束三原则中,A 原则(照应词原则)是针对照应词或反身代词说的,由于汉语没有反身代词,就无法检验这个原则。汉

语中的“自己”接近照应词,使用情况比较复杂。我们以下面的一个支配范畴例子为基础来考察支配范畴在研究中的意义,进一步认识结构关系的意义。这个范畴为:

NP1 V 自己

我们最后将提出一个传递性原则来说明汉语“自己”的活动规律,并证明支配范畴对于解释“自己”的活动规律仍然是一种重要的概念。考虑具体例子:

张三_i 认识自己_i
 我_i 认识自己_i
 他_i 认识自己_i
 你_i 认识自己_i 吗?

但在多层支配范畴中,汉语的“自己”不受照应原则限制。也就是说,“自己”还可以在支配范畴之外受到约束:

张三_i 说[李四_j 了解自己_{i,j}]
 张三_i 说[李四_j 喜欢自己_{i,j}]
 张三_i 说[李四_j 原谅了自己_{i,j}]
 张三_i 说[李四_j 批评了自己_{i,j}]

张三_i 不喜欢[李四_j 批评自己_{i,j}]
 张三_i 不喜欢[李四_j 对自己_{i,j} 的批评]

黄正德(1982)已经注意到这类现象。Chomsky(1986b)年提到“长距离约束”(long-distance binding)的概念,也谈到了英语中的长距离约束问题。黄运骅(1984)、Tang(1989)等提出了“阻

断效应”,认为汉语“自己”受根句(matrix clause,即主句)主语约束的条件之一是和该句所有主语在人称特征(包括数量、性别)上保持一致,否则只在小句中受约束。我们认为阻断效应有比较强的解释力,但还会遇到一些问题,下面我们来解释一些反例,并给出更一般的规则。

问题的本质是“自己”的先行词从内层支配范畴向外层支配范畴传递的问题。宋作艳(2001)发现,先行词的指称向后的传递并不受性别限制,或者说“自己”受前面先行词约束时并不受先行词的性别限制:

他_i知道[她_j会批评自己_i]_j

由于汉语口语中没有性别的区分,性别的限制在口语中更不存在。至于数量限制,宋作艳(2001)还发现,由于“自己”是指单数,要求先行词也是单数,如果是复数,往往都有“都”等单位把复数分解成单数:

张三_i知道[他们_j都会批评自己_i]_j

张三_i知道[他们_j有的会批评自己_i]_j

他们_i都知道[李四_j会批评自己_i]_j

这里的“他们”由于受“都、有的”的限制,表示“他们”中的每一个人。由此可以看出,先行词从内层支配范畴向外层支配范畴的传递和人称的数、性并没有直接关系,只和人称有关系。宋作艳(2001)认为长距离约束只与第三人称有关系,它要求所有主语必须都是第三人称(包括第三人称代词和指人的指称语)。我们同意这种看法。下面我们的分析将扩展到第一、第二人称的例子,并认为三种人称可以从理论上统一起来。

就我们下面的材料看,从根本上说,长距离约束只和单数人

称的一致性有关系。即只要人称不一致,就不可以传递:

我_i知道[她_j会批评自己_i]
 我_i知道[他_j会批评自己_i]
 我_i知道[你_j会批评自己_i]
 你_j知道[她_k会批评自己_j]
 你_j知道[他_k会批评自己_j]
 你_j知道[我_k会批评自己_j]
 他_k知道[你_l会批评自己_k]
 他_k知道[我_l会批评自己_k]
 她_k知道[你_l会批评自己_k]
 她_k知道[我_l会批评自己_k]

如果人称一致,先行词就可以传递,不限于第三人称:

我_i知道[我_i会批评自己_i]
 你_j知道[你_j会批评自己_j]
 张三_i知道[李四_j会批评自己_i]
 张三_i知道[他_j会批评自己_i]
 张三_i知道[她_j会批评自己_i]
 他_j知道[李四_k会批评自己_j]
 她_j知道[李四_k会批评自己_j]

 她_k知道[她_k会批评自己_k]
 她_k知道[他_l会批评自己_k]
 他_l知道[她_l会批评自己_l]

由于第一人称、第二人称情况简单,后面重点讨论第三人称的情况。

徐烈炯(1993)讨论了“自己”的长距离约束,提出了一个例子来反对阻断效应:

总统_i 请我_j 坐在自己_i 的旁边

其实类似的例子我们还可以举出很多:

张三_i 让我_j 站在自己_i 的后面

张三_i 让我_j 站在自己_j 的门口

张三_i 把我_j 放在自己的床上_i

张三_i 被我_j 拦在自己_i 的门前

张三_i 要我_j 选举自己_i

张三_i 请我_j 选举自己_i

这些都是人称不同,但先行词可以长距离约束的。阻断效应在这里确实遇到了问题。从汉语的情况看,在多个支配范畴内,“自己”前面的名词只要是指称人的题元,都可能约束“自己”。至于这种约束到底能否实现,要取决于词义、认知或语境意义等条件。比如上面的第一例,“我”不可能站在自己的后面,所以“我”不能约束“自己”,“自己”就只能被“张三”约束,这是认知条件决定的。如果把“后面”换成“门口”,“自己”就可以既被“张三”约束,也被“我”约束,这就是第二句的情况。其他几句都是有歧义的。徐烈炯的例子也属于认知的情况,因为“我”不可能坐在自己的旁边,所以“自己”只能受“总统”约束。

现在来概括汉语的约束规则。对于那些真正的多层支配范畴,先行词的传递可以一直进行下去,只要不出现记忆上的障碍。下面是三个支配范畴的嵌套:

{小王_k 说(张三_i 知道[李四_j 会批评自己_{i/j/k}]))}
 {小王_k 说(张三_i 知道[李四_j 会选举自己_{i/j/k}]))}

满足上述条件而不参与约束的一些先行词,也是因为词义、认知、语境等条件的限制:

小王_k 说[张三_i 知道[李四_j 会看望自己_{j/k}]]
 小王_k 说[张三_i 知道[李四_j 会拜访自己_{j/k}]]
 小王_k 说[张三_i 知道[李四_j 会送自己_{j/k}一本书]]

“看望、拜访、送”这样的外向动词,从认知上说一般定不会把动作指向自己,所以内层支配中的主语李四不会约束“自己”。

一般地说,在内层支配范畴内,“自己”前面的指称人的题元,无论是宾语或主语,都可能充当先行词。如果是跃出内层支配范畴,就只有主语能充当先行词:

小王_k 告诉张三_i 说[李四_j 会批评自己_{i/k}]
 小张_k 告诉小王_i 说[张三_j 知道[李四_j 会选举自己_{i/j/k}]]
 小张_k 告诉小王_i 说[张三_j 知道[李四_j 给自己_{i/j/k} 留了一个位置]]

黄正德(1982)已经注意到“自己”有选择主语作先行词的倾向,比如:

老王_i 告诉小李_j 说自己_i 要来
 老王_i 送给小李_j 一本自己_i 的书
 老王_i 不愿意同小李_j 谈自己_i

徐烈炯(1992)举出两个反例:

他_j一直被我们_i当作自己_i的榜样
 小张_j把小李_i关在了自己_i的屋里

根据我们对支配范畴的修正定义,以上的例子都是可以解释的。“自己”前面的指人的题元从原则上看都可能参与约束“自己”。前一个例子的“他”之所以没有约束“自己”,是因为“他”不能拿“自己”作榜样,这是认知条件决定的。徐烈炯的第二个例子,实际上“自己”可以受“小张”和“小李”约束。黄正德的第一个实例中,“小李”是外层支配范畴中的一个宾语,所以不约束“自己”。黄正德的第二个实例是单纯支配范畴,“自己”前面的“小李”、“小张”都有资格约束“自己”,但从认知上看如果书是小李自己的,并不需要送书,所以只能是“老王”约束“自己”。黄正德的第三个实例也是单纯支配范畴,但“同小李谈自己”这个片断本身就决定了不是小李谈自己,所以只能是“老王”约束“自己”。

根据以上分析,正如我们前面提到的,内层支配范畴以内并不区别主语和宾语,只要在“自己”前,都可能约束“自己”。在外层支配范畴中,先行词只在主语之间传递。

黄正德(1982)关于主语倾向性的观点,还暗含了一个结论,即“自己”可能越过内层支配到外层支配中寻找先行词,而根据我们以上的分析,“自己”首先要在内层支配范畴中获得先行词,如果内层支配范畴中位于“自己”前面的题元词不能作先行词,总是有词义、认知、语境的条件。

以上分析可以概括成先行词传递原则:

对于[S2 [S1 [S0]]]这样一种支配范畴的嵌套关系:

1. 在内层支配范畴 S0 内,“自己”首先受前面指称人的题元词约束(支配范畴的核心可以是谓语复杂形式)
2. 如果内层支配范畴中的先于“自己”的题元词和外层支

配范畴的主语在人称上一致,则“自己”的约束关系可以向外层支配范畴的主语传递,但不向宾语传递

3. 在可能的先行词中,“自己”最后被哪一个约束,取决于具体的词义、认知意义和语境

从以上讨论可以看出,汉语中“自己”的受约束方式和英语的 anaphor 不完全相同,但都需要通过支配范畴或通过对支配范畴的扩展来进行规则描述,而和规则有出入的例外,都可以通过语义、认知和语境等得到解释。更概括地说,汉语的名词、代词、指称词和“自己”等的约束规则基本上都需要通过支配范畴来描述,这就进一步证明了支配范畴的必要性。至于描述出来的规则是否和英语相同,显然有语言系统或语言参数的差异。前面我们讨论过支配要涉及核心和补足语两个概念,所以支配关系本质上是句法结构关系,现在我们可以说句法结构关系是约束关系的必要概念。正是因为结构关系是根本性的,即使汉语中的长距离约束能够越出支配的几个层次,这种约束仍然要受到主语和宾语概念的限制。而主语和宾语的概念正是以动词为核心形成的句法关系概念。

以上分析说明我们不应该纠缠约束原则是否在汉语中有效,因为每个语言都有每个语言的规则,而应该透过具体语言的规则来认识人类语言的初始概念和普遍原则。这涉及 Chomsky(1981)所说的原则与参数问题(principles and parameters)。普遍原则(universal principles)对世界上所有的儿童来说都是普遍的,但不同的语言有不同的语法结构,儿童需要学习这些特别的结构,这就是参数(parameters)。下面这些描写都可以看成参数:

1. 汉语的一般疑问句并不需要把助动词挪动到前面。
2. 汉语的疑问代词不需要移动。

3. 汉语的定语在动词前面。

参数反映了不同语言的差异,但是参数是有限制的,体现了一定的规则。比如:

What would you like

What are you going to do

Where are you going

Wh-类词总是转换到句首。不可能是有的转换到句首,有的不转换。从原则和参数的角度看,我们更应该观察支配范畴这一初始概念(或者更进一步说是结构关系这一初始概念)是否对汉语研究有价值。这个问题本质上又是句法结构关系和语义结构关系的相互制约问题。

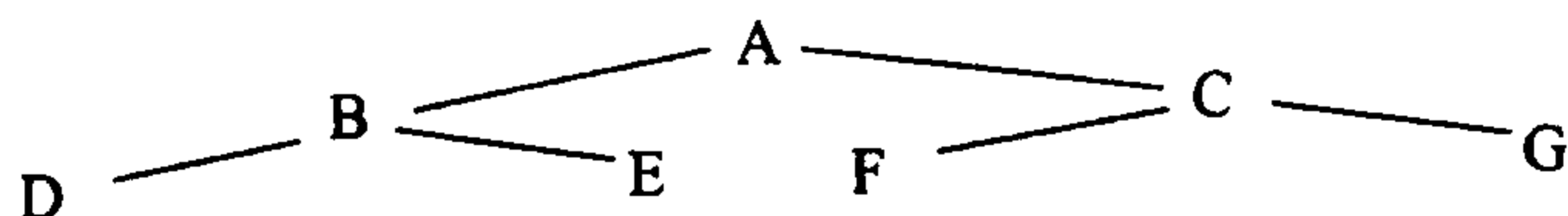
12.3 支配与结构关系

通过前面的英语和汉语实例分析可以说明,Chomsky (1981)的支配理论在判定题元角色、说明移动条件等方面确实具有相当大的解释力。下面试图进一步说明,支配概念的解释力是因为重新考虑结构关系而获得的,更早的理论模型 c-command 由于不考虑结构关系,解释力受到限制。

12.3.1 c-command

20 世纪 60 年代开始出现的 c-command 理论(Klima, 1965; Langacker, 1969; Lasnik, 1976; Reinhart, 1976; Stowell, 1981; Aoun and Sportiche, 1981),试图说明指称、转换和短语结构内部构造的关系。c-command 即 constituent-command,可以翻译成成分控制。要理解 c-command,首先要理解统辖(Dominate)。

考虑下面的图形：



以上节点统辖(Dominate)情况是：

A 统辖 B、C、D、E、F、G

B 统辖 D、E

C 统辖 F、G

所谓 c-控制(c-command), 根据 Chomsky(1995, p35)的定义是：

α c-ommands β if α does not dominate β and every γ that dominates α dominates β .

这一定义极其严密抽象。Radford(1988)曾经给过一个更为通俗的定义：

X c-commands Y iff (= if and only if) the first branching node dominating X dominates Y, and X does not dominate Y, nor Y dominate X (a branching node is a node which branches into two or more immediate constituents) 。

根据以上定义, 上面树形图中的控制关系有：

B c-控制 C、F、G

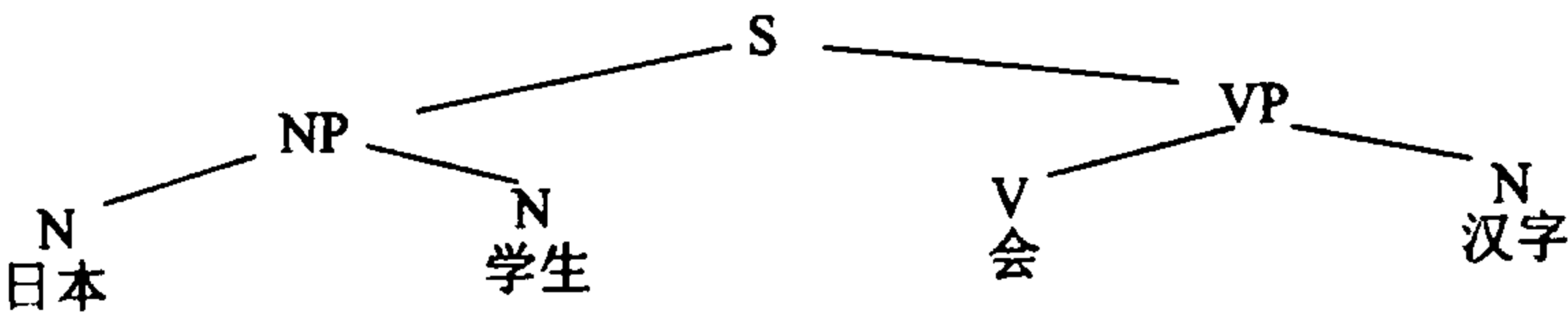
C c-控制 B、D、E

D c-控制 E

- E c-控制 D
- F c-控制 G
- G c-控制 F

以上 c-command 中,B 成分控制 C、F、G,但 B 的成分 D、E 并不成分控制 C、F、G,有时候为了强调这种区别,我们把 B 成分控制 C、F、G 称为直接成分控制或直接 c-控制。在有些语言的句法属性中,D、E 成分控制 C、F、G 的关系也是有价值的,我们把这种成分控制称为间接成分控制。一般情况的成分控制都是指直接成分控制。Chomsky 定义中的“every γ that dominates α dominates β ”规定了 D、E 不能控制 C,因为统辖 D、E 的 B 并不统辖 C。

考虑一个汉语的实例:



$\langle S(NP[N \text{ 日本}][N \text{ 学生}])(VP[V \text{ 会}][N \text{ 汉字}])\rangle$

直接成分控制或通常所说的成分控制包括下面的情况:

施控成分	受控成分
日本学生	会汉字,会,汉字
会汉字	日本学生,日本,学生

容易看出,NP 在成分上控制 VP 及其所辖成分 V 和 N,VP 也在成分上控制 NP 及其所辖的两个 N。

间接成分控制包括下面的情况:

施控成分	受控成分
日本	学生,会汉字,会,汉字
学生	日本,会汉字,会,汉字
会	汉字,日本学生,日本,学生
汉字	会,日本学生,日本,学生

一般地说,对于无论多长的语类序列,比如下面的语类序列:

$[_x[_y\text{ ABCDEF...}][_z\text{ GHIJKLM...}]]$

我们有:

- Y 直接控制 Z 及其 Z 统辖的所有成分。
- Z 直接控制 Y 及其 Y 统辖的所有成分。
- Y 统辖的每个成分间接控制 Z 及其 Z 统辖的所有成分。
- Z 统辖的每个成分间接控制 Y 以及 Y 统辖的所有成分。

对于汉语的一个多次嵌套实例来说:

$([_s[_{NP}\text{ 祖父的祖父的祖父}][_{VP}\text{ 是父亲的父亲的父亲的父亲的父亲的父亲}]])$

我们可以有:

- NP 直接控制 VP 及其所统辖的成分
- VP 直接控制 NP 及其所统辖的成分
- NP 所统领的每个成分都间接控制 VP 及其所统辖的所有

成分

VP 所统领的每个成分都间接控制 NP 及其所统辖的所有成分

这里的要点在于,由于一个姊妹成分 A 可以 c-command 另一个姊妹成分 B 及其所统辖的所有成分,A 就可以 c-command 很远的成分。比如:

[_{NP} 汪锋][_{VP} 断定邵永海去山东了]

[_{NP} 汪锋][_{VP} 断定郭锐知道邵永海去山东了]

[_{NP} 汪锋][_{VP} 断定孔江平认为郭锐知道邵永海去山东了]

这里的 NP 不仅 c-command 后面的 VP,而且 c-command 由 VP 统辖的全部成分,当然也包括距离很远的“山东”。下面我们会看到,这种不考虑句法结构关系的 c-command,最终会出现问题。

一般文献中提到的 c-command,实际都是直接成分控制。在英语中,用直接成分控制可以解释部分先行词和回指词(anaphor)的指称规则,即回指词要求有一个在指称上一致的先行成分,并且这个先行成分必须成分控制回指词。Radford(1988, p117)把这一规则总结为:

An anaphor (回指词) must have a appropriate c-commanding antecedent.

(回指词必须有一个在指称上一致的对回指词直接成分控制的先行成分)

Radford (1988, p117)给出的下面两个句子中的第 2 个句子违背了上述条件,所以是不合法的:

The soldiers disgraced themselves (disgrace: 丢脸)

$\langle_s (\langle_{np} [\text{d The}] [\text{n soldiers}]) (\langle_{vp} [\text{v disgraced}] [\text{n themselves}]) \rangle) \rangle$

$\langle_s (\langle_{np} [\text{The} [\text{n soldiers}']] [\text{n behaviour}]) (\langle_{vp} [\text{v disgraced}] [\text{n themselves}]) \rangle) \rangle$

在第一个句子中,回指词 themselves 和先行成分 the soldiers 的指称是一致的,并且 the soldiers 在成分上 c-command 回指词 themselves,所以这个句子是合法的。在第二个句子中,The soldiers' behaviour 在成分上控制 themselves,但 The soldiers' behaviour 和 themselves 的指称不一致;The soldiers 和 themselves 一致,但在成分上不直接控制 themselves。

在汉语中,上面回指规则的限制似乎不起作用:

[那些士兵的行为][羞辱了自己]

这里可以提出两个解释。一个解释是,汉语的“自己”并不是严格意义上的回指词。另一个解释是,汉语中的回指条件可以宽松到间接成分控制,因为“那些士兵”和“自己”实际上是间接成分控制关系。比较上面汉语和英语的差别可以看出,英语中的 anaphor 要被先行成分直接控制,间接控制不符合语法,而汉语中的所谓回指词可以间接控制。

尽管汉语的规则和英语不同,但可以看出,直接成分控制在人类的某些语言中是起作用的。问题是,即使在英语中,直接成分控制所起的作用是否是充分的?如果我们把上面的英语实例扩展一下,就有:

The students knew that the soldiers disgraced themselves

如果按照成分控制的定义, the students 也成分控制后面的 themselves, 但很明显, the students 的指称和 themselves 并不相同。可见, A 控制 B 并不等于 B 和 A 一定有同指关系, 即施控和受控之间并不一定有同指关系。这里的关键问题在于, 由于 c-command 没有考虑围绕某个中心语法关系, 一个成分可以控制距离很远的另一个成分。比如上面的实例还可以延长:

[_{NP} The teachers] [_{VP} heard that the students knew
that the soldiers disgraced themselves]

根据 c-command 的定义, NP(the teachers)能够 c-command 后面的 VP 以及 VP 中的所有成分, 但事实上 the teacher 和 themselves 并不同指。从理论上说, 这里 VP 的长度还可以无限扩展, NP 也可以无限远的 c-command 后面的成分, 建立在这种 c-command 基础上的约束关系都会出现问题。

如果我们再把上面的英语实例改动一下:

[The soldiers in the battle] [disgraced themselves]

根据前面所给出的关于成分控制的定义, the soldiers 在这里并不成分控制后面的 themselves, 但 the soldiers 和 themselves 显然有同指。这又说明, 不在成分控制范围内的施控和受控之间不一定没有同指关系。

我们认为成分控制之所以控制不住反身代词的指称, 根本原因在于成分控制只考虑成分和层次, 没有考虑语类序列的句法结构关系。下面进一步来证明这一点。

12.3.2 m-command

Anoun 和 Sportiche (1983) 提出了 m-command 的概念。m-command 即 maximal-command, 也可以翻译成最大控制。这里的“最大”指的是最大投射。Chomsky (1995, P35) 在定义 c-command 的基础上定义过 m-command:

α c-ommands β if α does not dominate β and every γ that dominates α dominates β . Where γ is restricted to maximal projections, we say that α m-commands β

Ouhalla (1999, p194) 给出了一个更为通俗的定义:

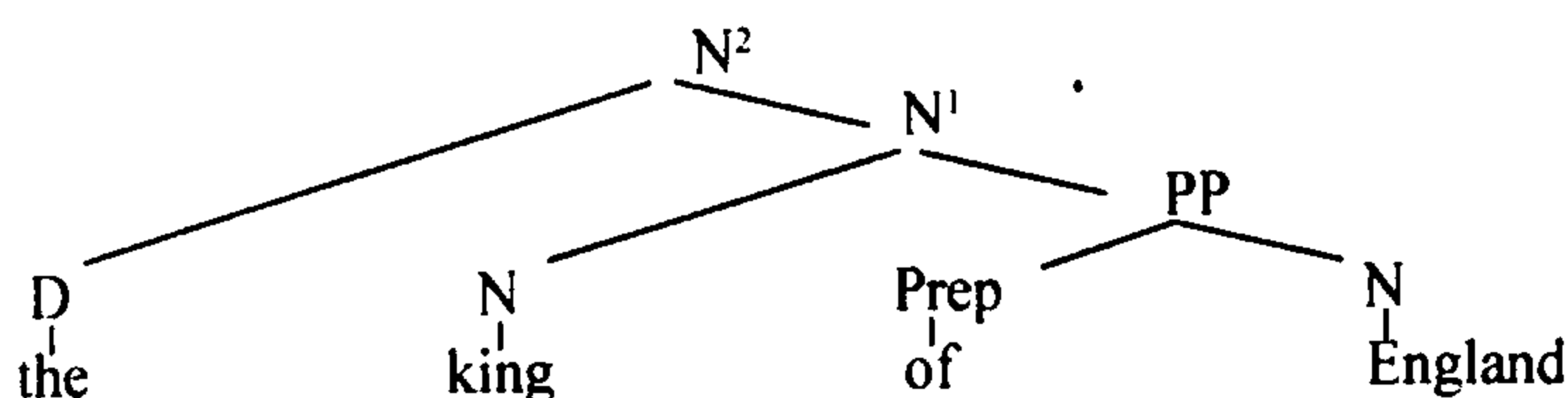
α m-ommands β iff:

- (i) α does not dominate β and β does not dominate α
- (ii) The first maximal projection dominating α also dominates β

m-command 这一概念的实质是, 凡是统辖 (dominate) α 的每个最大投射 (maximal projection) 同时也统辖 β , 并且 β 不统辖 (dominate) α , 于是 α 就在最大投射中控制 (m-command) β 。容易看出, α 和 β 是同一个最大投射下的成分。我们先以一个名词短语类说明什么叫最大投射。在下面的名词短语中:

the king of England

$\langle_{N^2} (\text{Spec } the) (\text{N}^1 [\text{N king}] [\text{pp of England}]) \rangle$



这个短语中的 king 是核心,是最小投射(minimal project), N^1 是间接投射(intermediate project), N^2 是最大投射(maximal projection), N^2 所统辖(dominate)的成分中,除了 N^1 统辖 king,其他成分都不统辖 king,于是 king 就 m-控制 N^2 中除了 N^1 以外的所有成分,包括[the]、[of England]、[of]、[England]。通过 king 和 the 的关系可以看出 m-command 和 c-command 的主要区别,king 可以 m-command 前面的 the,但不 c-command 前面的 the。

用 m-command 来解释反身代词的控制问题,就更有效。再考虑上面的实例:

(The students) ([knew] [that the soldiers disgraced themselves])

尽管 the students 按照定义成分控制后面的 themselves,但由于 the students 和 themselves 不在一个最大投射中,即 the students 并没有最大控制后面的 themselves,所以不会和 themselves 同指。另一方面:

[The soldiers in the battle] [disgraced themselves]

尽管 the soldiers 按照定义没有成分控制后面的 themselves,但由于 m-command 后面的 themselves,所以 the soldiers 和 themselves 同指。

为什么 m-command 有效? 我们认为根本原因在于 m-com-

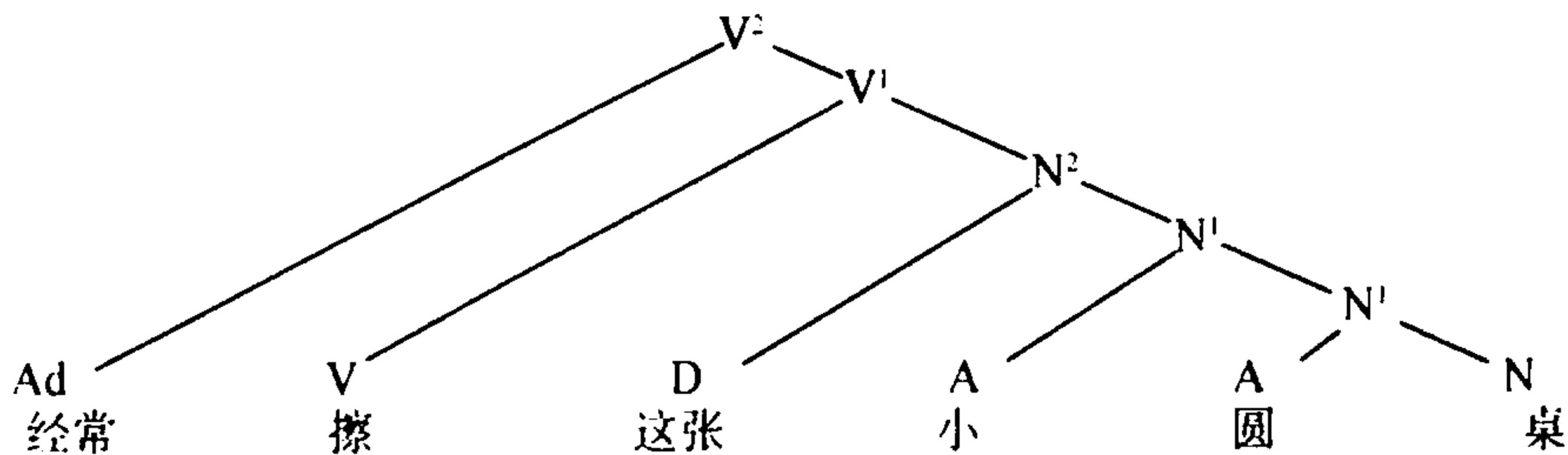
mand 中实际已经有了结构关系的概念。m-command 的定义要有最大投射的概念,由于最大投射必须涉及核心(head)的概念,即最大投射是关于某个核心的最大投射,所以 α m-command β 已经断定 α 和 β 共同处在一个以某个核心为中心的 maximum projection 中, α 可以是核心,也可以不是核心。核心和附加是一种重要的结构关系。概括地说,以上分析可以看出 m-command 和 c-command 的区别:

1. m-command 是最大投射条件下的 c-command,所以控制范围有了最大投射的限制。

2. 由于引进了最大投射的概念,而最大投射又必须涉及核心(head)的概念,于是最大投射内部之间的各个成分不再是平等的分布关系,而是中心和补足语的关系,从这种意义上说,m-command 已经有了句法结构关系的概念,而 c-command 没有。

从成分的角度看,m-command 是指最大投射下成分和成分的关系,但不是结构和成分的关系。考虑较复杂的汉语实例:

〈经常(擦[这张小圆桌])〉



这里有两个最大投射,一个是 VP(V²)最大投射,一个是 NP(N²)最大投射。在 VP 中,m-command 关系有:

α		β
经常	m-command	擦这张小圆桌、这张小圆桌、 小圆桌、圆桌、擦、桌
擦这张小圆桌	m-command	经常
擦	m-command	经常、这张小圆桌、小圆桌、圆 桌、桌
这张小圆桌	m-command	经常、擦

容易看出,如果不包括斜体字,正好构成 c-command 的情况。
在 $NP(N^2)$ 中, m-command 的关系有:

α		β
这张	m-command	小圆桌、圆桌、小、圆、桌
小	m-command	圆桌、圆、桌、 这张
圆	m-command	小、桌、这张
桌	m-command	圆、 小
圆桌	m-command	小、这张

同样,如果不包括斜体字,正好构成 c-command 的情况。从以上实例分析可以看出 m-command 把控制关系限制在了最大投射中,但在同一个最大投射中, α 能够 m-command 的成分比 α 能够 c-command 的成分多。比如,作为一个根本性的区别,在最大投射 VP 中,“擦”能够 m-command“经常”,但不能 c-command“经常”。与此类似,在 NP 中,根据 c-command 定义,“桌”是不 c-command“小”的,但是从关系上看,“小”和“圆”实际上都

是“桌”的补足语,都是修饰“桌”的,因此以“桌”为核心,构成了一个最大投射,“桌”能够 m-command“小”。

12.3.3 支配的本质和作用

A m-command B 并没有限定 A 是核心。Chomsky 的支配概念则明确限定了核心,如果 A 支配 B, A 就是核心,这样做的目的是要更进一步通过支配关系解决题元的约束、移位、控制等问题。所以要解决句法研究的很多问题,最终还是要回到支配概念上来。下面要论证,支配本质上仍然是对句法结构关系的恢复,尤其是对早期支配关系的恢复。

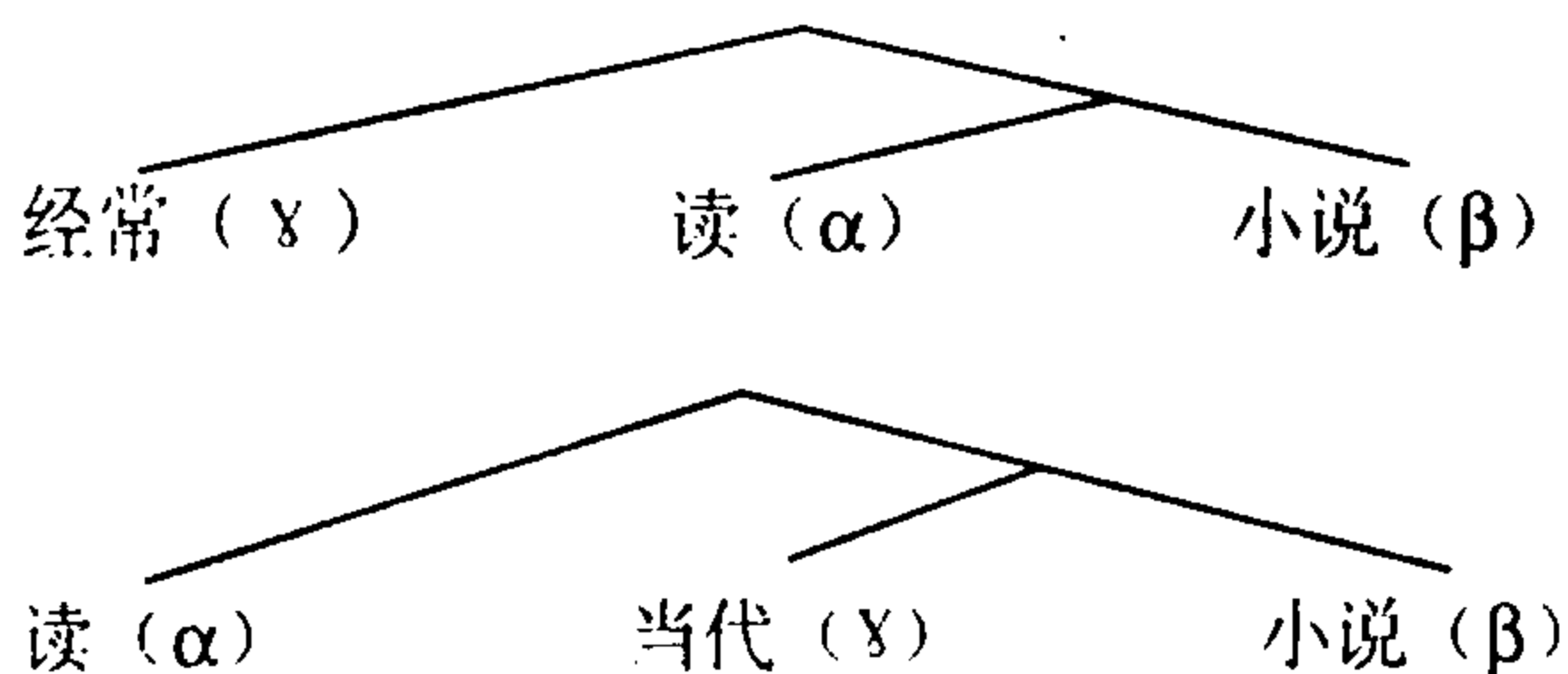
支配有好几种定义。但离不开核心的概念。核心和附加的关系是支配的本质。首先看建立在 c-command 上的定义 (Chomsky, 1981, p163)。

α governs β if and only if

(i) $\alpha = X^0$

(ii) α c-commands β and if γ c-commands β then γ either c-commands α or is c-commanded by β

$\alpha = X^0$ 表示 α 一定是核心成分。根据这里的定义, γ 在 α 和 β 之间的位置有下面两种典型情况。



用线性方式表示就是:

$$\begin{aligned} &(\gamma)([\alpha][\beta]) \\ &[\alpha][\gamma\beta] \end{aligned}$$

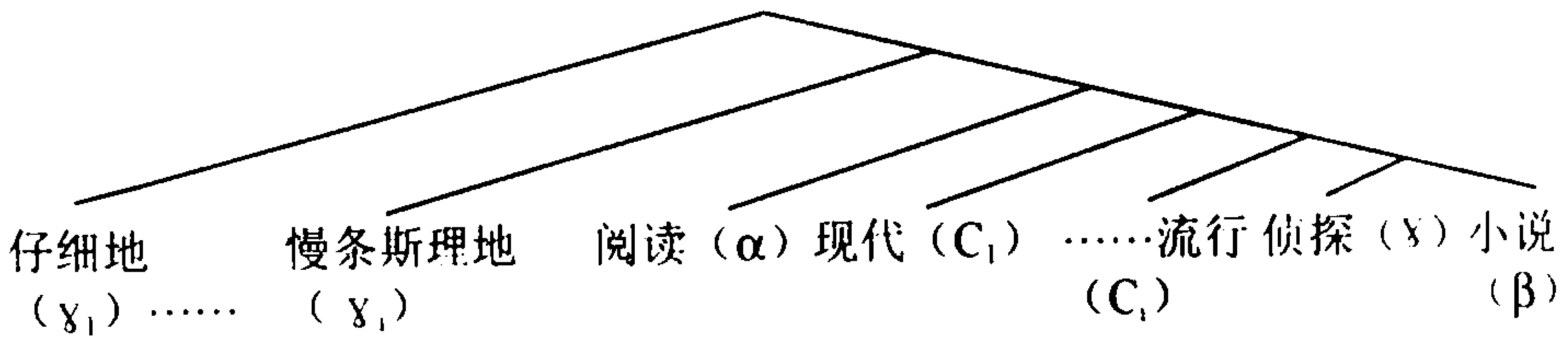
容易看出，“读”不仅支配“小说”，也支配“当代”。

满足 γ 条件的成分可以多个出现，其他成分(C_i)也可以出现，就有：

$$\begin{aligned} &\{\gamma_1(\langle \gamma_2 \cdots (\gamma_n) \rangle)([\alpha \beta]) \\ &(\alpha)(C_1 C_2 \cdots C_n[\gamma \beta]) \end{aligned}$$

比如：

仔细地(γ_1)…… 慢条斯理地(γ_i) 阅读(α) 现代(C_1)……
流行(C_i) 侦探(γ)小说(β)



从这个实例可以看出，如果依赖 c-command 定义的支配概念，“阅读(α)”可以支配很远的“小说(β)”，不过这时候核心“阅读”和附加成分“小说”的结构关系是不会变的。根据上述定义，也可以看出，“阅读(α)”实际上也支配“侦探(γ)”，这时“阅读(α)”和“侦探(γ)”并没有核心和附件的关系。

有了 m-command 的概念，在一个最大投射中核心支配附加成分的概念基本上已经形成。Chomsky(1986, p8)给出了在 m-command 基础上的支配定义：

α governs β if and only if α m-commands β and every barrier for β dominates α

这里的 barrier(语障)基本上可以理解成一个最大投射。我们前面说过, m-command 本身需要以最大投射为前提, 这里再引入 barrier, 等于再一次以最大投射为前提, 不过并不是重复。这里的本质含义是: 从 α 出发, α 和 β 在一个最大投射中, 从 β 出发, α 和 β 也在同一个最大投射中。考虑前面提到的实例:

(_{VP} 经常擦 [_{NP} 这张小圆桌])

“擦”支配“这张小圆桌”, 满足 Chomsky 的定义, 但“擦”不支配“这张”, 因为“这张”的语障是 NP, 这个 NP 并不统辖“擦”

Ouhalla(1999, p194)给出了一个更容易理解的关于支配的定义:

α governs β if and only if

(i) $\alpha = X^0$

(ii) α m-commands β

(iii) Minimality is respected

第 iii 个条件中的 Minimality 指的是 Minimality Condition, 即阻止一个核心去支配另一个核心构成的支配域中的成分。在这一点上和 barrier(语障)的限制是等价的。

既然在最大投射中引入了核心的概念, 那么核心对附加语就有一种支配关系。用 c-command(成分控制)定义这种支配显然存在问题。考虑下面实例:

(刚[来北京])

这里的“刚”尽管成分控制“来北京”, 但并不是核心, 所以不支配

“来北京”。“来”尽管是核心,但不成分控制“刚”,所以也不能支配“刚”。但是以“来”为中心,“刚、北京”为补足语,显然构成了一个最大投射。再看一个实例:

(小[方桌])

这里的“小”尽管成分控制(c-command)“方桌”,但并不是核心,所以不支配“方桌”。“桌”尽管是核心,但不成分控制“小”,所以也不能支配“小”。很明显,“桌”是核心,“方、小”都是“桌”的补足语,“小方桌”构成一个最大投射。

前面所考虑的 X 杠理论,其中涉及的根本属性之一就是最大投射问题。比较下面两个结构:

〈喜欢{这间(大[房子])}〉

〈我的{这间(大[房子])}〉

“喜欢”和“我的”所控制(c-command)的成分完全是一样的,也就是说,c-command 的分析结果相同。从结构语言学的直接成分理论看,以上两个实例的直接成分可以对应起来。但是,“我的这间大房子”只包含一个最大投射,而“喜欢这间大房子”不只包含一个最大投射。这里的根本区别就是两个直接成分的结构关系问题。如果不考虑结构关系,仅仅靠 c-command,结构中的最大投射反映不出来。如果最大投射反映不出来,则不能有效说明很多成分的移动或转换的机制。

前面讨论的 m-command,如果用来定义 government,以上的“来”就可以支配“刚”,“桌”就可以支配“小”,这是因为 m-command 实际上已经涉及了最大投射,而要定义最大投射,必须涉及核心(head)和支配,因为最大投射是围绕某个核心的最大投射,没有核心就无所谓最大投射。

支配的思想 Bloomfield(1933, 12. 8, p192)讨论过。对于下面的情况:

I know
watch me
beside me

Bloomfield 把这种选择关系称为支配(government), 其中 know 支配 I, watch 支配 me, beside 支配 me。Bloomfield (1933, p192)说:

This type of selection is called government; the accompanying form(know, watch, beside) is said to govern (or to demand or to take) the selected form (I or me).

Chomsky(1981, 2. 3, p50)对 government 作了解释, 我们认为其基本精神和 Bloomfield 是一致的, 主要区别在于对核心的处理不一样:

Case assignment is closely related to government, a crucial notion to which we will return to Chapter 3. Normally, Case is assigned to an NP by a category that governs it. As for notion "government", for the moment let us simply assume that the potential governors are the categories $[\pm N, \pm V]$ and INFL, that a category governs its complements in a construction of which it is the head (e. g. , V governs its complement in VP, etc.), and that INFL governs the sentence subject when it is tensed.

Chomsky(1981,p5)指出,支配理论的中心概念是结构的核心(head)和依赖核心的范畴的关系:

The central notion of government theory is the relation between the head of a construction and categories dependent on it.

关于支配的定义相当抽象,对支配概念的认识也在不断变化,但基本思想是稳定的,下面根据 Chomsky(1981,2.3,p50)的实例来进一步分析 government 的实质:

(8) John thought that he left his book on the table

Chomsky (1981,p50)解释说:

The matrix verb *think* governs its complement S^1 , but not any element (e. g. *he*) inside S^1 . The embeded verb governs its complements *his book* and *on the table*, but does not governs any element (e. g. , *his* and *table*) with in these categories. Thus, *his book* and *book* receive objective Case (the later, by percolation). The two occurrences of INFL govern *John* and *he*, assigning them nominative Case. The preposition *on* governs and assigns objective Case to its complement the *table*.

根据这里的解释,以上实例的支配关系有:

支配者	受支配者
think	he left his book on the table
left	his book, on the table
INFL(thought)	John
INFL(left)	he
on	table

以上 Chomsky 的支配关系和直接成分关系有下面的关系：

C ₁	C ₂	是否支配	是否是直接成分关系
INFL	John	支配	(直接成分)
think	that he left his book on the table	支配	直接成分关系
think	he	非支配	非直接成分关系
leave	his book	支配	直接成分关系
leave	on the table	支配	直接成分关系
INFL	he	支配	(直接成分)
leave	his	非支配	非直接成分
leave	table	非支配	非直接成分关系
on	table	支配	直接成分关系

根据以上比较可以看出,除了三种情况需要解释,以上的支配关系都属于姊妹成分关系。带括号的姊妹关系有一定的疑问,那是因为 INFL 这样的屈折成分是 GB 理论中深层结构的抽象成

分,所以不容易从经验上判定成分关系。可以说以上支配关系都是姊妹成分关系,但姊妹成分关系未必是支配关系,因为姊妹成分关系不考虑核心和附加的概念,而支配必须考虑核心和附加的概念。概括地说,由于支配理论引进了核心成分和附加成分(complement)的概念,从根本上说是引进了结构关系的概念,即引进了中心和补语这种结构关系。

前面我们还讨论过 Bloomfield 的支配概念。从以上对比还可以看出,支配理论在三个方面发展了 Bloomfield 的支配概念,这就是上面提到的需要解释的三种情况。下面来考虑这三种情况。

第一种发展是提出 INFL 支配主语的思想。这是支配理论为了解释主语的被支配问题而提出的一种假说,INFL 也是虚拟单位,是动词形态和情态的抽象概括。尽管 Chomsky 把 INFL 和主语 NP 分析成树形图上的姊妹关系,但由于 INFL 不是一个可观察到的特定的单位,而是若干可观察单位的抽象,INFL 和主语的关系已经不完全是线性关系。

第二种发展体现在 Chomsky 的以下认识中:

left 支配 his book

left 支配 on the table

从这种意义上讲,government 和直接成分在理解层次观念上不完全相同。比较下面直接成分理论和支配理论的分析结果:

直接成分分析结果:

$(_{vp}[_{vp} \text{left his book}] \text{ on the table})$

支配理论分析结果:

$(_{v_1} \text{left}[_{np} \text{his book}] [_{pp} \text{on the table}])$

第三种发展是重新定义直接成分的核心。以上的支配关系，按照 Bloomfield 的观点，可以分成向心结构和离心结构两类：

向心结构(斜体字为核心)	离心结构
to <i>think</i> that ...	John thought that...
to <i>leave</i> his book on the table	on the table

支配理论则把两类结构都看成统一起来，都有一个核心。比如，以上两个离心结构在支配理论中，核心分别是 INFL 和 on。根据支配理论，汉语中的“的”字结构的核心是“的”。这就涉及核心的经验判定的问题。

可以说，支配(government)实际上也是从直接成分理论发展出来的。支配成分就是短语的中心，是单个的词或语法形式，包括动词、名词、形容词、介词、形态等。受支配成分是支配成分的补语或与支配成分结合的另一个直接成分。比如下面两个句子：

- (1) speak the language
- (2) speak about the language

在(1)中，speak(governing)支配 the language，但并不是支配 language，而只是成分控制(c-command) language。在(2)中，speak 支配 about the language，但并不支配 the language，而只是成分控制 the language。the language 受 about 的支配。

正是因为支配和直接成分的理论渊源关系，在大多数情况下，我们也可以从直接成分的姊妹关系和核心的角度定义支配(government)的概念：

- 在 AB 中，如果：
1. A 和 B 互为姊妹成分

2. A 是 AB 的核心

我们就说 A 支配 B

被支配的补足语 (complement) 也可以是 trace。government 的概念对确定 trace 和 PRO 很关键。Chomsky (1981, p60) 提出了一个原则来区分 trace 和 PRO:

(13) If α is an empty category, then α is PRO if and only if α is ungoverned (equivalently, α is trace if and only if α is governed)

Chomsky (1981, 3.2, p50) 认为 a category governs its complements, 但 complement 的判定也不是完全没有问题。因此判定 complement 又是确定 trace 和 PRO 的关键。后来 Chomsky 等又不断对支配的定义有修正, 但基本精神和上面的讨论一致, 就是核心支配补足成分。

government 是 GB 理论的核心, 其重要性可以概括为两个最为重要的方面, 实际上都和识别题元有根本关系:

1. 约束 (binding) 要由支配范畴 (governing category) 来定义

约束的三个条件前面已经有讨论。Chomsky 最终是通过支配范畴而不是成分控制 (c-command) 来控制住约束的, 即约束只能用支配范畴来定义, 不能用成分控制来定义。由于支配本质上是一种结构关系, 这说明约束必须依赖结构关系。

2. trace 和 PRO 的区分要由支配来确定

Chomsky (1981, p60) 进一步提出了一个原则来区分 trace 和 PRO:

(13) If α is an empty category, then α is PRO if and

only if α is ungoverned (equivalently, α is trace if and only if α is governed)

Chomsky(1981, p51)指出了支配(government)的重要性:

As the example indicates, government rather than subcategorization is the relevant notion for Case-assignment-in this case, at S-structure. But government is also the relevant notion for subcategorization in D-structure in case (2I), and generally. Thus the theories of subcategorization, θ -marking and Case all fall within the general theory of government, at least in their essentials.

从更高层次看,这些问题又涉及语障(barrier)的问题,即线性片断有多长的情况下可以有支配关系和移位的可能,以上分析说明,无论是支配范畴还是语障,这些概念和范畴都离不开以某个成分为核心的结构,最终都要追问结构关系,追问功能关系。

12.4 句法结构关系的初始性

支配与约束理论由于技术上的复杂和术语的更新,不容易让人看出问题的实质。仔细分析 Chomsky(1981)的支配约束理论,实际上是重新启用句法结构关系来解决一系列重要理论问题。更具体地说,支配约束理论的核心问题之一是要解决语义解释中的根本问题,即怎么样从表层的语类序列判定语法角色和语义角色(论元角色),尤其是要解决空语类的论元角色问题,因为空语类是隐性的,不容易识别。早期的 c-command 为什么不能充分解决空语类问题? 主要就是因为 c-command 只有语类和直接成分的概念,没有结构关系的概念。支配约束理

论从根本上说是在充分考虑语义结构关系的基础上识别题元角色,即判定空语类的性质。人类语言到底需要哪些句法结构关系来描述,现在仍然不清楚。但支配概念至少证明了偏正结构关系的重要性。现在的问题是句法结构关系是否是初始的。让我们再来考虑 Chomsky (1965) 的取消句法结构关系的依据。Chomsky 给出过下面的实例:

sincerity may frighten the boy

要说明这个句子,Chomsky 认为传统语法涉及三种信息:

1. 词类信息
2. 句子成分信息,主宾语问题,即功能概念
3. 名词的再分类信息:可数名词、抽象名词、动物名词等;

动词的再分类信息:及物、不及物等

句子成分的词类问题在短语结构规则中得到了描写,但成分的功能并没有得到直接描写。Chomsky 认为,尽管句子成分等概念和范畴是不同的,但仍然可以通过重写规则来描写。这种处理方式基于两个理由(Chomsky, 1965, p69—70):

1. function 是冗余信息,不必专门处理。
2. 类似下面的例子用主语、宾语等功能概念来处理是不可能的:

(7) (a) John was persuaded by Bill to leave

(b) John was persuaded by Bill to be examined

(c) what disturbed John was being regarded as
incompetent

Chomsky(1965,p71)讨论了这个问题。所谓主语,就是相对特定句子的 NP,所谓宾语,就是相对于 VP 的 NP。而这些信息都包含在由短语结构构成的理论中,即树形图本身已经有这些信息。更明确地说,由于短语结构规则或树形图提供了范畴和层次的概念,因此也就提供了主语、宾语等句子成分概念。按照 Chomsky (1965,p70)的定义:

主语:[NP,S] (即主语是句子的 NP)
 谓语:[VP,S] (即谓语是句子的 VP)
 直接宾语:[NP,VP] (即宾语是 VP 的 NP)
 主语动词:[V,VP] (即述语是 VP 的 V)

这样功能的信息都已经在重写规则中了。

我们曾经讨论过结构功能原则(陈保亚,1999,2.1.3),认为结构关系是初始概念,并不是冗余信息。有些结构关系不可能通过语类推导出来,其中我们讨论的实例中就涉及核心和附加的推导问题。比如:

塑料玻璃
 电脑桌子
 香蕉芒果
 香蕉苹果
 衣服帽子
 学生家长
 孩子父亲

这里两个直接成分的语类都是 N,整体结构的语类是 NP,但有偏正关系和并列关系两种,也就是说,相同的语类关系并不能断

定必然有相同的结构关系。

那么句法结构关系是否可以从语义结构关系推导出来呢？一般来说，动词都处在动词短语的中心，在汉语这样的语言中，多数情况下动词也是句子的中心。是否可以说动词就是中心呢？如果是这样，句法结构关系中的核心和附加成分的关系就可以通过动词和语义格的关系定义。但回答是否定的。再回顾前面讨论过的一些例子：

狗追猫

追猫狗

狗追猫

猫被狗追

是狗追猫

是追猫狗

以上每个实例中语义结构关系不变，句法结构关系却不同，从这种差别可以认识到句法结构和语义结构的本质区别。句法结构是表达层面的，语义结构是认知层面的，两者都和功能有关系。形式语法试图通过语类来解决功能问题，前面各章我们对一些主要模型和概念的分析实际上已经显示，要从语类推导功能有一定的困难。再考虑下面实例：

炒菜、煮菜、炖菜、烧菜、蒸菜

这里只有一种语义结构关系，即行为和受事，但有两种结构关系：如果是偏正结构，“炒”是核心；如果是述宾结构，“菜”是核心。可见句法结构关系和语义结构关系是不同层面的关系。

另一方面，短语并不只是动词短语或由动词转换出来的短语，还有名词短语、形容词短语、副词短语、量词短语等，这些短

语的核心都不是动词中心论能够解决的。

功能的概念可以从三个层面来理解：句法功能、语用功能、语义功能。本章重点讨论了语义功能和句法功能，以及这两种功能的关系。句法功能即传统语法中语言成分在结构关系中形成的功能，包括主语、谓语、宾语、定语、状语、补语等。语义结构功能即前面我们讨论过的语义格功能。

现在我们可以对句法的层面做一个归纳。句子可以分成下面几个层面：

	猫	吃了	老鼠
语类层面	NP1	V	NP2
语法功能层面	主语	谓语动词	宾语
语义功能层面	施事	行为	受事
语用层面	话题	述题	

下面是成分标注，可以看出在转换过程中各个层面的角色不是一一对应的：

1. 主语 $[_N \text{猫}^{\text{施事}}]_{\text{话题}}$ 吃了 宾语 $[_N \text{老鼠}^{\text{受事}}]$
2. 主语？ $[_N \text{老鼠}^{\text{受事}}]_{\text{话题}}$ 主语？ $[_N \text{猫}^{\text{施事}}]$ 吃了
3. 主语？ $[_N \text{老鼠}^{\text{受事}}]_{\text{话题}}$ 被 宾语 $[_N \text{猫}^{\text{施事}}]$ 吃了
4. 主语 $[_N \text{猫}^{\text{施事}}]_{\text{话题}}$ 把 宾语 $[_N \text{老鼠}^{\text{受事}}]$ 吃了

这里名词性成分被括号标记，左括号的右下方表示这个成分的语类，左括号的左上方表示这个成分的句法功能，右括号的右下方表示这个成分的语用功能，右括号的左上方表示这个成分的论元功能。在这些语义结构相同的转换式中，语类和论元没有变化，语法功能和语用功能在变化。第2句中的“老鼠”的地位各家有不同的解释，由此引出了汉语有没有主语以及主语和话题的关系问题。标注中后面带问号的是有争议的功能标注。我

们认为争论的实质在于怎样定义主语和话题。从线性标记看,汉语中主语和话题没有差异,但从生成句子的角度看,第2句显然不只是通过初始连接和替换生成的,还需要叠加程序或转换程序,前面相关章节已经有讨论。

现在我们回到一个重要问题上来,即传统的句子成分和直接成分的关系。从前面的分析可以看出,支配、最大投射、X 杠理论这样一些重要的而且相互联系的概念实际上是对传统的中心词分析和直接成分分析的继承和发展。以汉语为例(没有精确到 IP):

$$\{ s(\text{NP 猫})(\text{VP}[\text{V 吃}][\text{NP 老鼠}]) \}$$

这里的“吃”是核心,“猫”和“老鼠”围绕着“吃”构成一种支配关系,也构成一个最大投射,构成一个 S-bar,同时也构成一个句子(S),这正是中心词分析的本质。但中心词分析不考虑直接成分的概念,而这里的支配、最大投射、X 杠都承认直接成分关系。生成语法对中心词分析和直接成分分析的扩展在于把支配、最大投射、X 杠扩展到了 VP、NP、AP、AdP、PP 等各种语类或句子,等于承认了层次和核心是所有组合的基本属性。前面我们从初始概念的角度论证了这一点。前面我们讨论过,符号的线性组合不是单位的线性组合,而是直接成分的线性组合,现在我们可以进一步说,符号的组合不仅仅是直接成分的语类线性组合,而且是具有一种结构关系的语类线性组合。为了更进一步理解这里的实质,我们再考虑前面提到的汉语实例:

他经常擦这张小圆桌

$$[s \text{ 他} [\text{VP 经常} [\text{VP 擦} [\text{NP 这张} [\text{NP 小} [\text{NP 圆} [\text{N 桌}]]]]]]]]]$$

这是一个我们前面提到的右线性结构,我们可以从右边的“桌”

开始不断地叠加成分,形成一种直接成分的线性组合。但是仅仅认识直接成分线性组合还没有完全认识线性组合的性质,因为当叠加到“这张”的时候,叠加的结构还是 NP,还是一个以“桌”为中心的 NP 最大投射^①。但当叠加到“擦”的时候,性质就发生了根本的变化,转入以“擦”为核心的 VP 投射。这里的支配关系发生了变化,也就是结构关系发生了变化,这种变化会对约束、移位、控制、赋格等产生不同的作用。

从上面的例子看好像这种支配关系或结构关系的变化是语类的变化,但情况不总是这样。比如:

喜欢学习汉语

[_{VP} 喜欢 [_{VP} 学习汉语]]

在叠加的过程中,VP 并没有改变,但中心从“学习”转移到了“喜欢”。

其实,正是因为 c-command 只考虑层次而不考虑围绕中心成分所构成的域,所以没有办法解决大量空语类问题和同指问题。一个先行词和回指词的距离可以远到什么范围? 一个成分可以移动到多远? 都会遇到困难。后来由于支配、m-command 等概念考虑了中心成分所构成的域,这些问题才得到解决,界限(bounding)、邻接(subjacency)、语障(barrier)等概念才可以得到比较明确的限定,约束理论、界限理论、控制理论等才走向成熟。而围绕中心成分所构成的域,就是句法结构关系的域。甚至从根本上说,X-bar 理论、投射理论都离不开核心和依附的句法结构关系概念。

到此我们可以看出,结构关系中最为重要的关系是支配和

^① 如果把“这张”看成是“这张小圆桌”的中心,会有不同的分析,但不影响这里要讨论的问题。

受配关系,我们前面讨论的“塑料玻璃”有歧义,不能通过语类推导结构关系,就和支配受配有关系,即是理解成两个支配还是一个支配一个受配的问题。因此说结构关系是初始概念,至少支配关系是一个关键证据。至于传统语法中的其他句子成分,如主语、宾语,是否可以通过语类推导出来,还可以进一步研究。由于语类和次范畴化是可以在词库中标注的,除此以外的结构关系是不可以在词库中标注的,因此从语类和次范畴化能够推导出哪些结构关系,是如何推导出来的,是语言认识论的重要问题。

思考问题

1. 支配概念的定义存在哪些问题?
2. 能否通过语类模式来预测结构关系?

结语：从程序到条件

回到 Saussure 思考的根本问题上来。Saussure 认为符号的组合是线性组合，结构语言学试图在这个基础上发展出发现程序，并仅仅靠替换和分布获得单位和规则，类推出组合的语类，并通过语类来解释各种功能。

我们的分析说明，自然语言在符号层面不存在单位的线性组合，即使那些看上去和前后都可以发生组合的单位，也都是在以直接成分的方式和结构关系进行线性组合：

{(the[man])comes}

有时候好像是两可的，那是因为有歧义：

三个研究室的老师

→三个[研究室的老师]

→[三个研究室的]老师

在每一个既定意义中，都是直接成分按照一定结构关系的线性组合。

我们还进一步论证了转换的必要性。自然语言并不完全符合线性原则，有些言语片断是通过在线性原则基础上变形得到的，这不是一个描写的优选问题，而是描写是否充分的问题。比如：

将军+他→将他的军

我(知道[他去过上海了]了)→我知道他去过上海了

因此,初始连接、替换、叠加等线性程序并不能生成人类语言的全部句子,转换也是生成人类语言句子的必要程序之一。不过,线性组合仍然是自然语言最基础的组合,连接、替换、叠加也是必要的生成程序,而且语言的递归性是靠这些线性程序获得的。

结构语言学的线性发现程序曾经试图在不依赖语素、词的条件下,从分布、替换来获得音系单位,然后获得语素,再通过语素的分布类获得话语。我们的分析说明,音系层面的单位和规则的获得需要语素、词的条件,这不仅仅是生成语法所说的优化问题,而且是必要条件。比如:如果只从语音出发,汉语普通话双音节调型的分布是有空格的:

55+55	55+35	55+214	55+51
35+55	35+35	35+214	35+51
214+55	214+35		214+51
51+55	51+35	51+214	51+51

仅仅靠语音层面的信息,我们只能断定 214+214(实际为 21+214)在表中不出现,并不知道上声会有哪些变体,214+214 跑到哪个变调组合类型中去了。只有预先获得了语素,才可能判定 214+214 变成 35+214。比如“土改、涂改”的例子。

但是,这不是说发现程序目标有问题,而是发现程序的方法有问题。在田野调查中,发现程序仍然是一个不可绕过的问题。如果我们先获得语素,在语素的基础上提取核心语音单位和语音规则,并通过平行周遍条件获得语法层面上的规则单位和规则,发现程序仍然是可行的。

我们还论证了功能研究的必要性。无论是线性组合规则还是转换规则,都需要考虑语法结构关系、语义结构关系和语用结

构关系,考虑关系就是在考虑功能,因此功能的概念是语法研究的初始概念,完全从语类分布来认识深层结构或线性结构是不充分的。c-command 之所以不能充分解释论元的移动、约束和控制的条件,就是因为 c-command 只考虑语类而不考虑结构关系,支配概念最终能够对论元的移动、约束和控制的条件作出比较好的说明,就在于支配概念最低限度的考虑了核心和附加这一最为初始的结构关系。

结构语言学试图还原单位,提出了有价值的分布、替换线性程序,是对语言性质的重要认识;生成语法提出了从有限单位和规则生成无限合法句子的目标,也是对语言性质的重要认识。但研究线性程序、转换程序以及音系和句法的各个层面,都更多地在研究句子的还原和生成程序。目前对生成程序研究的工作已经相当深入了,对提取单位和规则的还原程序也有一定的研究。但是,从还原方面看,Saussure 谈到的获得单位的困难仍然没有解决;从生成方面看,Chomsky 生成合法句子的要求也远远没有达到。我们认为最根本的问题在于我们对各种还原和生成程序的平行周遍条件还很不清楚。比如“X 鞋”:

布鞋、棉鞋、金鞋、冰鞋、跑鞋……

这些是规则组合还是不规则组合?是两个规则单位还是一个规则单位?在给出平行周遍条件以前,可能生成这样的不合法片断:

· 走鞋、· 滑鞋

如果我们把平行条件限制在 X 表示质料和性质,“X 鞋”就是一种周遍关系,就可以生成任意多的组合(还要受韵律条件限制):

钢鞋、铝鞋、玻璃鞋、木头鞋、钢化玻璃鞋……

可见提取单位和规则与理解和生成是同一个问题的两个向度，即还原程序和生成程序是同一个问题的两个方面。这两个方面必须要考虑平行周遍条件。这个问题不仅存在于线性组合中，也存在于转换中，比如本书引言中提到的复合动趋势宾语的转换现象：

后置：	我刚才从冰箱里	我刚才从柜子里
	拿出来一瓶酒	拿出来一件衣服
中置：	我刚才从冰箱里	我刚才从柜子里
	拿出一瓶酒来	拿出一件衣服来
前置：	？我刚才从冰箱	？我刚才从柜子
	里拿一瓶酒出来	里拿一件衣服出来

为什么这里的宾语“一瓶酒、一件衣服”移动到“出来”前面就不太合法？需要研究平行周遍条件，即这里的前置移动必须满足未然条件，由于句子中已经有“刚才”，就不允许有这种移动。至于我们是用线性程序、构式还是转换程序来分析这里的移动过程，已经不是最迫切的工作。又比如我们在书中提到的一个学界经常讨论的实例：

[纸张粉碎][机]

要对该实例进行减缩还原是很困难的，即不能还原出“纸碎机”。也就是说，这个实例不是扩展替换生成的。如果我们把不能用减缩替换还原的结构看成是转换的结构，上述实例的生成过程就是：

[粉碎纸张]+[机]→[纸张粉碎][机]

即“粉碎纸张”和“机”组合时,由于受到韵律的影响,出现转换,宾语移动到动词的前面。这是建立在替换基础上的转换。我们也可以认为“纸张粉碎机”是在基式上叠加形成的:

[粉碎]+[机]→([粉碎][机])

(纸张)+([粉碎][机])→(纸张)([粉碎][机])

这时就不需要转换。其实这两种生成机制都可以,涉及整个语法体系的处理方案,即到底线性组合是连接、替换生成的,还是连接、叠加生成的。但是这里更迫切的工作是,是什么条件激发了“纸张粉碎机”和“碎纸机”的语序差异。

在汉语句子的线性组合中,哪些 VN 只有述宾结构(吃饭),哪些兼有述宾和偏正结构(炒饭),平行周遍条件还没有找到。汉语“把”字句的转换程序可以有好几种,但哪些动词可以启动“把”字句的转换规则,平行周遍条件并不清楚。自然语言中大量的线性组合规则和转换规则都存在类似的问题,这已经成了生成合法句子的瓶颈问题。有些生成程序已经相当有深度了,但平行条件的周遍性问题却没有得到充分的研究。考虑一个显著的实例。结构语言学和早期生成语法在解释下面一类名词短语时,认为名词 king 或 kings 是核心(head)。比较:

this king of England

these king of England

that king of England

those kings of England

生成语法后期的 DP(determinal phrase)认为限定词(determin

er)those 是核心。后来的 dP 理论注意到句子(IP)在转换成名词短语时的平行关系：

John translated the book→John's translation of the book
John invited Jack→John's invitation of Jack

一些学者认为不仅动词的核心是屈折变化形态 I, 名词短语的核心也是屈折变化标记, 比如性、数、格的标记。这一理论能够比较好的解释名词和指定语的一致关系。比如 those kings of English, 这时复数 s 是核心, king of English 是补足语(Complement), 由于核心是复数 s, 所以选择复数形式 those 作指定语(Spec)。概括地说, 这个形态成分 d 以 NP 为补足语, 共同构成 D'。D' 以限定语为指定语, 共同构成 DP。甚至还有不同的分析。这样的分析已经相当精巧了, 无论我们是否同意这些分析, 它们对名词短语的内在结构的认识相当有深度。但无论这样的认识有多深, 都是在解释生成程序而不是生成条件。更为根本的问题是, 在英语中, 从句子到自指的转换, 哪些动词能够转换成 V-ion 形式, 哪些动词不能, 即下面格式的转换条件是什么：

NP1-V-NP2→NP's-V-ion of NP2

John translated the book→John's translation of the book

John invited Jack→John's invitation of Jack

John composed the music→John's composition of the music

John explained the music→John's explanation of the music

.....

能否找出周遍条件, 直接影响到这一格式生成力的强弱。如果能够找出周遍条件, 就是一种生成规则, 如果找不到周遍条件,

就只能是一种理解规则。至于把名词的形态理解成名词短语的核心还是把名词或限制语理解成名词短语的核心,这些选择并不能解决平行周遍问题。激发转换的条件是迫切要解决的问题,这需要做大量的经验研究工作而不是符号处理工作。

概括地说,连接、替换、叠加和转换几种程序的生成能力是相当强大的,这些程序所能生成的序列远远超出了自然语言中的话语序列,必须考虑这些程序的平行周遍条件,才能生成合法的自然语言句子。我们认为第二语言学习在言说和写作上出现大量错误,机器翻译乃至自然语言理解未能取得大的进展,都和我们还没有找出话语还原和生成的平行周遍条件有很大的关系。将来的研究应该有一个大的转向,把更多的精力投放到经验工作中,以便寻找平行周遍条件,即从还原程序和生成程序的研究转向还原条件和生成条件的研究。和还原程序和生成程序相比,还原条件和生成条件是生成合法句子的更为根本的问题。

在没有找到组合的平行周遍条件以前,各种复杂的现象就应该暂时存放在词库中。比如词库中就应该注明哪些 V 和 N 组合可以有两种结构关系(比如“炒菜”),哪些 V 可以有“把”字句的转换(比如“吃、喝”),其实实际工作中有不少学者已将开始这么做,只是没有充分论证这么做的理由。词汇功能语法把大量信息存放在词库中,是目前的研究的一个趋势。但是词汇功能语法并没有给出充分的理由,所以有些满足平行周遍条件的组合也都存放到词库中了。我认为言语片断存放在词库中的根本理由还是一个平行周遍条件的问题。平行不周遍的规则必须存放在词库中,才能避免生成不合法的句子:

一碗炒饭 / 一碗买饭

把问题解决了 / 把问题知道了

一旦我们找出了线性组合程序或转换程序的平行周遍条件,就

需要把原来存放在词库中的一组相似的组合删除,而在规则库中给一个记录。

这个过程就是不断从话语中还原规则和单位或固定格式,结构语言学的发现程序实际上就是在做这样的工作,但由于结构语言学只考虑从单位的分布来认识规则,而且不区别平行和周遍两个条件,所以发现程序并没有真正找到可以生成合法句子的单位,更没有说明单位赖以成立的规则。生成语法对规则进行了充分的研究,并且把很多问题上升到原则,解决了语言生成的程序问题。但由于仍然没有充分研究生成程序得以启动的平行周遍条件,生成语法要生成合法句子的目的也没有达到。看来我们的主要工作还是要回到规则和单位的还原上来,我们需要还原出平行规则和平行周遍规则,同时得到满足这两种条件的单位和结构,才能真正生成合法的句子。

这又把我们引向了一个最为根本的问题:什么叫语言能力?语言能力在理解和生成句子中有多大的作用?

如果语言能力是人不同于任何动物的天赋,那么语言能力不仅仅表现在能够通过有限的单位和规则理解和生成无限的句子,语言能力还包括通过有限的话语片断还原出可以用来理解并生成无限句子的有限规则和单位。语言能力是由还原和生成两个部分组成的,生成语法重点在研究生成部分,结构语言学发现程序重点在研究还原部分。要充分认识语言能力,必须从还原和生成两个向度展开研究。在这两个向度的研究中,如果不考虑规则类推的平行周遍条件,都会遇到困难。

考虑平行周遍条件必须考虑语音、句法、语义、认知、语用等各个层面的因素,因此要生成合法的句子,还必须研究语言系统和语言运用。我们认为 Chomsky 虽然区分了语言能力和语言运用两个概念,但没有区分语言能力和语言系统两个概念。Chomsky 以前的语言学家也没有自觉区分语言能力和语言系统。Saussure 尽管提出了语言系统的概念,没有提到语方能力

但我们认为,应该承认语言能力是由遗传决定的,这种能力主要表现为人类具有通过有限的单位和规则生成无限句子的能力,在一定的句子中还原出规则和单位的能力,以及通过任意性原则将概念和语音形式相结合的编码能力。人以外的动物由于不具备这样的能力,不可能学会严格意义上的语言。但人的语言能力只能说明人可能学会语言,不能说明自然语言的差异,也不能说明为什么具有同样遗传基因的种族在不同的社会环境中学会的是完全不同的语言结构和范畴。和语言能力不同,语言系统是由特定范畴组成的复杂系统,它是在历时过程中形成的,从历史的角度看,它体现了人类语言能力和后天经验的相互作用。语言运用不仅仅是语言能力的运用,也是语言系统的运用。Chomsky 后来提到的和原则相对的参数,实际上就是由语言系统决定。在语言活动中,语言能力、语言系统、语言运用(或言语)是三个不同的概念。要还原出单位和规则,生成合法的句子,还必须从这三个方面展开研究,还原能力和生成能力是先天的,具体还原出什么规则和单位则和后天经验有关系。由于要考虑经验事实,具体语言中平行周遍规则的还原就相当困难,需要展开大量的语料调查分析工作。

因此,要生成合法的言语片断,当前的真正瓶颈在于如何从大量语言材料中还原出单位、规则及其平行周遍条件。

符号解释

A	形容词;施事
Ad	副词
AGR	一致关系
AP	形容词短语
APd	形容词的字短语
Atr	分类性修饰语,属性修饰语
Comp	补足语;补足语标记
CP	Complementiser Phrase
D	区别词
GR	语法关系
IA	项目与配列
INFL;I	曲折成分
IP	Inflection Phrase
IP	项目与变化
L	量词
LF	逻辑形式
LIR	词汇插入规则
N	名词
NP	名词短语
PF	语音形式
PP	介词短语
Prep	介词
PosN	领属短语
PreN	代词作修饰语的领属短语

PS rule	短语结构规则
Q	数词
QLP	数量词短语
QP	量词短语
V	动词
VP	动词短语
VPd	动词的字短语
XPd	的字结构
* X	X 是不合法的句子
? X	X 是不太合法的句子

本书有关层次、语类、功能的体例格式采用严格的表示方法：

买这张小桌子

$$\{_{VP}^{Head} \langle_V \text{买} \rangle^{Comp} \langle_{N2}^{Spec} (QP \text{ 这张}) (N1^{Comp} [_{Adj} \text{小}]^{Head} [_{N} \text{桌子}])^{patient} \rangle \}$$

相同的直接成分用相同的括号。“小”和“桌子”互为直接成分，都用括号[]，“这张”和“小桌子”互为直接成分，都用括号()。更多的层次以此类推。不强调直接成分时，不用严格的表示方法。语类在左括号的右下方，句法结构功能在左括号的左上方，语义结构功能在右括号的左上方。在上面的例子中，整个短语的语类是一个 VP，用{_{VP}表示，整个短语的核心是“买”，用^{Head}⟨表示，该 VP 短语后面的“这张小桌子”的句法功能角色是补足语(Comp)，用^{Comp}⟨表示，该补足语的语类是 NP(N2)，用⟨_{N2}表示。该名词短语的语义功能角色受事(patient)，用^{patient}⟩表示。

这里的 NP 又是一个最大投射，也有层次、语类和功能。表示方式与上面类似。

不影响理解时，标记方式不一定严格遵守。

人名译名对照外文排序

以下译名对照只涉及本书参考书目中有中文和外文译名的情况。为便于检索参考书目和索引,译名对照分为中文排序和外文排序两种。除非是普遍公认的译名,本书提到的外国人名尽量用原文或英文,以避免译名分歧带来的检索错误,有分歧的译名在人名对照中尽量都给出。

Austin J. L.	奥斯汀
Bao Zhiming	包智明
Bloomfield L.	布龙菲尔德
Chao Y. R.	赵元任
Chen M. Y.	陈渊泉
Cheng C. C.	郑锦全
Chomsky N.	乔姆斯基
Church A.	丘奇
Dragunov A.	龙果夫
Duanmu San	端木三
Feng Shengli	冯胜利
Fillmore C.	菲尔墨
Fodor J. A.	福德
Frege F.	弗雷格
Halliday M.	韩礼德
Harris Z. S.	海里斯
Hartman L. M.	哈忒门
Hockett C. F.	霍凯特

Huang C. T. James	黄正德
Jespersen Otto	叶斯泊森
Katz J. J.	凯茨
Lakoff G.	拉科夫
Lapolla Randy	罗仁地
Leibniz G.	莱布尼茨
Moore G.	摩尔
Wittgenstein L.	维特根斯坦
Ross J.	罗斯
Russell B.	罗素
Ryle G.	赖尔
Sapir E.	萨丕尔
Saussure F.	索绪尔
Swadesh	斯瓦迪士
Sweet H.	斯维特
Tesnière L.	特思尼耶尔
Trubetzkoy N. S.	特鲁别茨科依
Wang William S-Y.	王士元
Wells R. S.	威尔斯
Xu L. J.	徐烈炯
Zhang Hongming	张洪明
Смерницкий А. И.	斯米尔尼兹基

人名译名对照汉语排序

奥斯汀	Austin J. L.
包智明	Bao Zhiming
布龙菲尔德	Bloomfield L.
陈渊泉	Chen M. Y.
端木三	Duanmu San
菲尔墨	Fillmore C.
福德	Fodor J. A.
冯胜利	Feng Shengli
弗雷格	Frege F.
哈忒门	Hartman L. M.
海里斯	Harris Z. S.
韩礼德	Halliday M.
黄正德	Huang C. T. James
霍凯特	Hockett C. F.
凯茨	Katz J. J.
拉科夫	Lakoff G.
莱布尼茨	Leibniz G.
赖尔	Ryle G.
龙果夫	Dragunov A.
罗仁地	Lapolla Randy
罗斯	Ross J.
罗素	Russell B.
摩尔	Moore G.
乔姆斯基	Chomsky N.

丘奇	Church A.
萨丕尔	Sapir E.
斯米尔尼兹基	Смерницкий А. И.
斯瓦迪士	Swadesh
斯维特	Sweet H.
索绪尔	Saussure F.
特鲁别茨科依	Trubetzkoy N. S.
特思尼耶尔	Tesnière L.
王士元	Wang William S-Y.
威尔斯	Wells R. S.
维特根斯坦	Wittgenstein L.
徐烈炯	Xu L. J.
叶斯泊森	Jespersen Otto
张洪明	Zhang Hongming
赵元任	Chao Y. R.
郑锦全	Cheng C. C.

主要概念译名对照外文排序

凡是已经有对照译名的,或者我们能够对译的,尽量给出。
有些术语是在汉语文本中使用的,或者是我们为了讨论问题的方便提出的,但还没有合适的西文对照,暂时空缺。

	当事
	广用式变调
	窄用式变调
accusative	宾格、受格
adjunct	附加语
adjunction	附加
agent	施事
anaphor	回指词、照应词
antecedent	先行成分,先行词
A-position	论元位置
argument	论元,主目
argument rank	论元层阶
argument structure	论元结构
artificial language	人工语言
atomic fact	原子事实
atomic proposition	原子命题
base	基础
base component	基础部分
base phrase-marker	基础短语标记
beneficiary	受益格

binding theory	约束理论
bound phrase	黏合式短语
bounding theory	边界理论
Case Filter	格过滤条件
Case Grammar	格语法
Case Theory	格理论
categorial feature	范畴特征
causative	使役格
c-command, constituent-command	c-控制, 成分控制
check	核查
cognitive	认知
combined phrase	组合式短语
command	控制
Comp; COMP	补足语标记, 补足语
complement	补足语
Complement Clauses	补语小句
Complementizer	补足语标记
Complementiser Phrase	补足短语
connotation	含义
constituent linearity	成分线性组合
context-free	上下文自由的
context-sensitive	上下文敏感的
control theory	控制理论
dative	与事
delete	删除
demonstrative pronoun	指示代词
denotation	指称
denoting	指称
derivation	派生

descriptive adequacy	描写的充分性
determiner	限定成分
Determiner Phrase	限定短语
dominate; domination	统领、统辖
DP	限定短语
epistemology of language	语言认识论
ergative	作格
ergative verb	作格动词
expansion	扩展
experiencer	感受格, 感事
Extended Standard Theory	标准理论
finite state grammar	有限状态语法
finite state language	有限状态语言
finite state machine	有限状态自动机
focus	焦点
formative	构词形式、语言形式
free alternation	自由交替
function	功能
Functional Grammar	功能语法
functional structure	功能结构
GB	支配与约束理论
Generalized Phrase Structure Grammar	广义短语结构语法
Generative Grammar	生成语法
Generative Semantics	生成语义学
Generative Transformational Grammar	转换生成语法
governing category	支配范畴
governing	支配

Governing and Binding Theory (GB)	支配与约束理论
government	支配
government theory	支配理论
government; govern	支配
government-binding	支配与约束
governed	受支配的
governor	支配者
grammatical category	语法范畴
grammatical function	语法功能、句法功能
grammatical relation	语法关系
head category	核心范畴
immediate c-command	直接成分控制
immediate constituent	直接成分
immediate constituent of similar generation	同辈直接成分
immediate constituent of unsimilar generation	不同辈直接成分
immediate constituent relation	直接成分关系
immediate domination	直接统领、直接统辖
INFL; I	屈折成分
inflection	屈折成分
Inflection Phrase	屈折短语
information focus	信息角度
initial	声母
initial state	开始状态
instrument	工具
intermediate c-command	间接成分控制
interpretive	解释性的
island constrain	孤岛限制

juncture	音渡
Katz-Postal Hypothesis	凯茨—波斯塔假说
kernel sentence	核心句
language competence	语言能力
language faculty	语言才能
language performance	语言运用
lexical decomposition	词项分解
lexical category	词汇范畴
lexical insertion rule	词汇插入规则
Lexical Phonology	词汇音系学
lexical rule	词汇规则
Lexical-Functional Grammar	词汇功能语法
lexicon	词库
light verb	轻动词
linearization	线性化
linguistic universals	语言普遍性
linguistic intuition	语言直觉
linguistic meaning	语言学意义
liquids	流音
logical form (LF)	逻辑形式
logical positivism	逻辑实证主义
long movement	长距离移动
maximal projection	最大投射
m-command; maximal-command	最大控制
meaning	意义
merge	合并
metagrammatical	元语法
middle structure	中动结构, 中动句
middle verbs	中动动词

minus morpheme	负语素
morph	语子
morpheme alternants	语素交替
morpheme; simple form	语素
morphophoneme	语素音位;形态音位
morphophonemics	语素音位学
move- α	α 移动
node	节点
nominative case	主格
non-overt anaphor	非显性回指词
NP-trace	名词语迹
O	受事
obligatory transformation	必选转换
oblique case	介词宾格
operator	算子
Ordinary coordination	并接原则
overt	显性的
overt anaphors	显性回指词
patient	受事
periphery	边际成分,卫星成分
periphery phoneme	边际音位
phoneme	音位
phoneme sequences	音位序列
phonemic level	音位平面
phonetic alternants	语音交替
phonetic form (PF)	语音形式
phonetic representation	语音表达式
phonology	音系学
Phrasal Verb	短语动词

phrase	短语
phrase structure rule	短语结构规则
phrase-marker	短语标记
phrase-words	短语词
predicate	谓语
presentation	表达式
principles and parameters	原则与参数问题
PRO	大代语
projection principle	投射原则
pronoun	代词
proper noun	专名
proposition	命题
quantifier	量词
reference	所指, 指称
reflexive	反身代词
Relational Grammar	关系语法
representation	表达式
rewrite rule	改写规则, 重写规则
selectional restrictions	选择限制
semantic feature	语义特征
semantic function	语义功能
semantic relation	语义结构关系
semantic representation	语义表达式
sign linearity	符号线性组合, 符号线条性
significant sequence	有意义的序列
signifier linearity	能指线条性
sisiter relationship of unsimilar generation	不同辈姊妹关系

sister relationship	姊妹关系
sister relationship of similar generation	同辈姊妹关系
slash category	斜线范畴
Spec	指定语
Specifier	指定语
S-structure	S-结构
Standard Theory	标准理论
Strict subcategorization rules	严格次范畴化规则
Structural Linguistics	结构语言学
structuralism	结构主义
subcategorization	次范畴化
subcategorization frame	次范畴化框架
subcategorization rules	次范畴化规则
subjacency	邻接
subjacency condition	邻接条件
subject	主语
substitution	替换
substitution rules	替换规则
symantic case	语义格
syntactic lexicon	语符库
syntactic relation	句法结构关系
syntactic unit	语符
tentative morpheme	暂拟语素
the minimalist program	最简方案
thematic hierarchy	题元层阶
thematic role	题元角色
thematic theory	题元理论
theta criterion	题元标准

theta-grid	格框架
topic	主题
trace	语迹
traditional grammar	传统语法
transformation	变换
transformation	转换
transformation cycle	转换循环
Transformational Grammar	转换语法
tree diagram	树形图
UG	普遍语法
underlying form	底层形式
unit linearity	单位线性组合
universal principles	普遍原则
utterance	话语
variable	变体
Verb Phrase Complements	动词短语补语
Verbal Complements	动词补语
well-formed	合式的
word-class	词类
X-Bar theory	X 杠理论; 语杠理论
zero morpheme	零语素
θ -position	题元位置

主要概念译名对照汉文排序

c-控制	c-command, constituent-command
S-结构	S-structure
X 杠理论;语杠理论	X-Bar theory
α 移动	move- α
必选转换	obligatory transformation
边际成分	periphery
边际音位	periphery phoneme
边界理论	bounding theory
变换	transformation
变体	variable
补足语标记	complementizer, COMP, C; Comp
标准理论	Extended Standard Theory
标准理论	Standard Theory
表达式	representation
宾格、受格	accusative
并接原则	Ordinary coordination
补语小句	Complement Clauses
补足短语	Complementiser Phrase
补足语	complement, Comp
补足语标记	Complementizer
不同辈直接成分	immediate constituent of unsimilar-generation
不同辈姊妹关系	sisiter relationship of unsimilar generation

核查	check
长距离移动	long movement
成分控制	c-command, constituent-command
成分线性组合	constituent linearity
重写规则	rewrite rule
传统语法	traditional grammar
词汇插入规则	lexical insertion rule
词汇范畴	lexical category
词汇功能语法	Lexical-Functional Grammar
词汇规则	lexical rule
词汇音系学	Lexical phonology
词库	lexicon
词类	word-class
词项分解	lexcial decomposition
次范畴化	subcategorization
次范畴化规则	subcategorization rules
次范畴化框架	subcategorization frame
大代语	PRO
代词	pronoun
单位线性组合	unit linearity
当事	
底层形式	underlying form
动词补语	Verbal Complements
动词短语补语	Verb Phrase Complements
短语	phrase
短语标记	phrase-marker
短语词	phrase-words
短语动词	Phrasal Verb
短语结构规则	phrase structure rule

反身代词	reflexive
范畴特征	categorial feature
非显性回指词	non-overt anaphor
符号线条性	sign linearity
符号线性组合	sign linearity
负语素	minus morpheme
附加	adjunction
附加语	adjunct
改写规则	rewrite rule
感受事,感事	experiencer
格过滤条件	Case Filter
格框架	theta-grid
格理论	Case Theory
格语法	Case Grammar
工具	instrument
功能	function
功能结构	functional structure
功能语法	Functional Grammar
构词形式、语言形式	formative
孤岛限制	island constrain
关系语法	Relational Grammar
广义短语结构语法	Generalized Phrase Structure Grammar
广用式变调	
含义	connotation
合并	merge
合式的	well-formed
核心范畴	head category
核心句	kernel sentence

话语	utterance
回指词	anaphor
基础	base
基础部分	base component
基础短语标记	base phrase-marker
间接成分控制	intermediate c-command
焦点	focus
节点	node
结构语言学	Structural Linguistics
结构主义	Structuralism
解释性的	interpretive
介词宾格	oblique case
句法功能	grammatical function
句法结构关系	syntactic relation
开始状态	initial state
凯茨-波斯塔假说	Katz-Postal Hypothesis
控制	command
控制理论	control theory
扩展	expansion
量词	quantifier
邻接	subjacency
邻接条件	subjacency condition
零语素	zero morpheme
流音	liquids
论元	argument
论元层阶	argument rank
论元结构	argument structure
论元位置	A-position
逻辑实证主义	logical positivism

逻辑形式	logical form (LF)
描写的充分性	descriptive adequacy
名词语迹	NP-trace
命题	proposition
能指线条性	signifier linearity
黏合式短语	bound phrase
派生	derivation
普遍语法	UG
普遍原则	universal principles
轻动词	light verb
屈折成分	INFL, I, inflection
屈折短语	Inflection Phrase
人工语言	artificial language
认知	cognitive
删除	delete
上下文敏感	context-sensitive
上下文自由	context-free
生成语法	Generative Grammar
生成语义学	Generative Semantics
声母	initial
施事	agent
使役格	causative
受格、宾格	accusative
受支配的	governed
受事	O, patient
受益格	beneficiary
树形图	tree diagram
算子	operator
所指	reference

题元标准	theta criterion
题元层阶	thematic hierarchy
题元角色	thematic role
题元理论	thematic theory
题元位置	θ -position
替换	substitution
替换规则	substitution rules
同辈直接成分	immediate constituent of similar generation
同辈姊妹关系	sister relationship of similar generation
统领、统辖	dominate; domination
投射原则	projection principle
卫星成分	periphery
谓语	predicate
先行成分、先行词	antecedent
显性的	overt
显性回指词	overt anaphors
线性化	linearization
限定成分	determiner
限定短语	Determiner Phrase
限定短语	DP
斜线范畴	slash category
信息角度	information focus
形态音位	morphoneme
选择限制	selectional restrictions
严格次范畴化规则	strict subcategorization rules
意义	meaning
音渡	juncture

音位	phoneme
音位平面	phonemic level
音位序列	phoneme sequences
音系学	phonology
有限状态语法	finite state grammar
有限状态语言	finite state language
有些状态自动机	finite state machine
有意义的序列	significant sequence
与事	dative
语法范畴	grammatical category
语法功能	grammatical function
语法关系	grammatical relation
语符	syntactic unit
语符库	syntactic lexicon
语迹	trace
语素	morpheme; simple form
语素交替	morpheme alternants
语素音位、形态音位	morphophoneme
语素音位学	morphophonemics
语言才能	language faculty
语言能力	language competence
语言普遍性	linguistic universals
语言认识论	epistemology of language
语言学意义	linguistic meaning
语言运用	language performance
语言直觉	linguistic intuition
语义表达式	semantic representation
语义格	symantic case
语义功能	semantic function

语义结构关系	semantic relation
语义特征	semantic feature
语音表达式	phonetic representation
语音交替	phonetic alternants
语音形式	phonetic form (PF)
语子	morph
元语法	metagrammatical
原则与参数问题	principles and parameters
原子命题	atomic proposition
原子事实	atomic fact
约束理论	binding theory
暂拟语素	tentative morpheme
窄用式变调	
照应词	anaphor
支配	governing
支配	government
支配	government; govern
支配范畴	governing category
支配理论	government theory
支配与约束	government-binding
支配与约束理论	GB
支配与约束理论	Governing and Binding Theory (GB)
支配者	governor
直接成分	immediate constituent
直接成分关系	immediate constituent relation
直接成分控制	immediate c-command
直接统领, 直接统辖	immediate domination
指称	denotation
指称	denoting

指称	reference
指定语	Spec, Specifier
指示代词	demonstrative pronoun
中动动词	middle verbs
中动结构、中动句	Middle structure
主格	nominative case
主目	argument
主题	topic
主语	subject
专名	proper noun
转换	transformation
转换生成语法	Generative Transformational Grammar
转换循环	transformation cycle
转换语法	Transformational Grammar
姊妹关系	sister relationship
自由交替	free alternation
组合式短语	combined phrase
最大控制	m-command, maximal-command
最大投射	maximal projection
最简方案	the minimalist program
作格	ergative
作格动词	ergative verb

索引

- adjunct 476, 480
Bloomfield 1, 7, 15, 32—34, 39, 48, 64, 103, 121, 122, 149, 163, 164, 170, 183, 186, 197, 213, 218, 246, 320, 331, 421, 458, 470, 475, 514, 546, 549, 550
c-command 87, 532, 533, 535, 537—540, 545, 551—553
Complementiser Phrase 462
determinal phrase 564
D 变韵 186
formative 39—42, 44, 47, 148, 170, 180, 181, 185
GB 理论 388, 391, 552
Harris 1, 13—16, 30, 35, 36, 38, 42, 44, 47, 73, 78, 101, 103, 121—124, 148, 161, 164, 165, 167, 169, 170, 187, 213, 215, 247, 285, 290, 305, 320, 437, 443, 458, 463—465, 470, 519
Katz-Postal 假说 399
m-command 538—540, 544, 546
Move- α 42, 109
X 杠理论 458, 460, 461, 465—470, 472, 476, 482, 484—486, 488, 490, 491, 493, 496, 507, 512, 514, 545
Z 变韵 186
本音 146, 155, 158, 161, 162, 170, 181, 182, 185, 187, 191, 192, 196, 199—202, 205, 240—244, 247
边界理论 524
变换 283
变换分析 305
变形规则 307
标准理论 2, 69, 283, 284, 293, 294, 301, 320, 321, 331—334, 399, 401, 404, 406—408, 413, 432, 436, 445, 454
表层结构 40, 180, 284, 300—302, 307, 319, 347,

- 364, 387, 407, 409, 434,
436, 438, 440, 445, 450,
452—455, 466, 468, 469,
477, 478
- 表达式 249, 277, 284, 407,
419, 420, 426—428, 430,
431, 439, 451
- 宾格 13
- 宾语 1, 103, 295, 315, 327,
331, 334—336, 346, 347,
364, 388, 427, 428, 437,
439, 446, 448, 529, 530,
553, 554, 556, 563, 564
- 搏动理论 218, 222, 223, 235
- 补足语 458, 461, 462, 467,
469, 471, 474—476, 478,
486, 487, 489, 512, 513,
531, 541, 545, 551
- 初始对立矩阵 260, 261, 265
- 初始连接 21, 24, 317
- 词汇插入规则 41, 296, 399,
401
- 词汇功能语法 303, 566
- 词汇假说 47
- 词汇音系学 177, 178
- 词项分解 401, 414, 419
- 词组本位 462
- 次范畴化 285—288, 290—
294, 297, 298, 300, 319,
330, 331, 333, 336, 337,
342—344, 349, 350, 363,
411, 440, 465, 478
- 递归性 25, 26, 36, 307, 308,
380, 382, 385, 458, 482—
489, 491, 494—497, 503,
506, 508, 512
- 叠加 21—29, 31, 32, 76—
78, 80, 83, 85, 92, 94, 95,
100, 107—109, 111—114,
119, 127, 129, 131, 133,
134, 137, 141, 143—145,
314, 488, 492, 498, 499,
558, 561, 564, 566
- 定式动词 462
- 动词补语 297—299
- 动词短语补语 297, 298
- 短语标记 40, 311, 344, 403,
443
- 短语结构语法 35, 81, 89,
96, 100, 101, 104—108,
110, 125, 134, 138, 144
- 对比 6, 8, 9, 11, 13—15, 25,
34, 51, 52, 55, 58, 60, 62,
63, 121, 166, 182, 183, 187,
188, 197, 200, 240, 252,
277, 281, 422, 549
- 多线性音系分析 196
- 发现程序 1—3, 14, 30, 147,

- 149, 156, 158, 162, 187,
192, 197, 331, 560
- 非自立 153, 161
- 非自主音系学 234
- 分析程序 14, 47, 130, 133,
145, 147, 150, 151, 179,
182, 185, 240
- 符号线条性 5
- 附加语 177, 486, 487, 489,
499, 545
- 改写规则 39, 41, 43, 44, 97,
101, 102, 105 — 107, 114,
127, 285, 290, 296, 309,
310, 320, 377, 381, 382,
407, 434, 460, 461, 482,
483, 485, 488, 490, 500
- 格语法 317, 319, 320, 331,
333 — 335, 340, 367, 399,
403, 404, 408 — 410, 414,
432, 435, 436, 446
- 构式 131, 132
- 孤岛限制 520
- 合并 20, 86, 94, 142, 175,
261, 285, 303, 306, 309 —
312, 346, 353
- 核心句 128, 129, 305
- 后结构语言学 2, 14, 153,
162, 187, 215, 251
- 还原替换 18
- 肌肉紧张理论 219
- 基础部分 80, 106, 284, 294,
296, 300, 320, 321, 330,
401, 444, 446
- 基式 15, 23, 26 — 29, 85,
106, 140, 143, 144, 470, 564
- 结构主义 14, 332
- 紧张理论 219
- 经事 351
- 句段关系 1, 5
- 控制理论 391
- 扩充 15, 18 — 24, 26 — 30,
32, 108, 114, 125, 127, 131,
141, 143
- 扩充替换 18
- 扩展的标准理论 301
- 离心结构 35, 458, 470, 514,
550
- 连接 6, 14, 15, 19 — 22, 24,
26, 29, 31, 32, 74, 77, 85,
87, 92, 96, 104, 107, 111,
112, 126, 127, 138, 145,
168, 170, 217 — 219, 222,
234, 237, 306, 308, 313,
314, 317, 318, 340, 341,
381, 401, 442, 443, 445,
557, 561
- 两层性 182, 184
- 量词 382, 400

- 论元 328, 336, 345, 351, 353, 360, 363, 366, 367, 377, 440, 520, 553, 557
- 论元结构 328
- 莫拉 229
- 内部论元 360, 363, 397
- 内破外破理论 223, 235
- 能量理论 223, 235
- 能指线条性 5
- 黏着语义格 358, 392
- 派生 34, 47, 68, 70, 71, 73, 102, 103, 126, 146, 148, 172, 175, 176, 180, 182, 186, 189, 190, 192, 204, 205, 240, 242, 243, 303, 308, 355, 377, 380, 407, 439
- 评价程序 2, 125, 384, 415
- 屈折成分 439, 462, 549
- 屈折短语 439, 462
- 认知 76, 110, 285, 338, 339, 341, 358, 360, 509, 528, 530, 531, 556
- 删除 127, 133, 258, 260, 261, 347, 360, 383, 443, 567
- 深层结构 15, 28, 72, 73, 80, 109, 140, 142, 284, 293, 300—303, 305—308, 313, 315, 317—320, 330, 331, 333, 334, 342, 377, 379—381, 399—403, 406—408, 410, 411, 413, 431, 432, 435, 436, 442, 443, 450, 451, 453, 455—457, 466, 468, 469, 478
- 生成语义学 302, 381, 399—402, 404, 407, 409, 410, 413, 415, 419, 420, 432, 435, 436, 446, 451, 454
- 施事 1, 128, 303, 318—320, 324, 327, 330—334, 339, 342, 344, 346, 351, 355, 357, 360—362, 365—367, 371, 388—390, 392—394, 396, 433, 438, 439, 556
- 实现规则 377
- 事件结构理论 368
- 受格 13
- 述语 554
- 树形图 43, 85, 90, 91, 103, 365, 400, 405, 406, 478, 533, 549, 554
- 双向唯一性原则 153—156, 158, 160, 161, 251
- 算子 400
- 题元理论 328
- 题元准则 329
- 替换 2, 6, 14, 15, 18, 19, 21, 22, 24—30, 32, 35, 38, 48,

- 50, 51, 54, 57, 62, 64, 78 — 80, 85, 92 — 94, 96 — 100, 106, 107, 109, 131, 132, 134, 138, 140, 143 — 145, 299, 306, 308, 313, 318, 348, 377, 381, 384, 385, 401, 409, 413, 426, 430, 434, 463, 464, 469, 479, 482 — 484, 486, 488, 491 — 494, 496 — 499, 502, 504, 505, 508 — 510, 512, 557, 563, 564
- 添加 133, 475, 487
- 投射原则 391, 436, 438, 440, 441
- 外部论元 363, 397
- 谓 语 103, 319, 331, 332, 338, 346, 364, 424, 426, 430, 516, 530, 554, 556
- 先行词 526 — 528, 530, 536
- 线条性 5, 6, 14, 26, 30, 81, 84, 96
- 线性结构 15, 28, 80, 108, 142, 303, 308, 310, 312, 313, 315, 333, 357, 367, 381, 519
- 响度阶 221, 223, 237
- 响度理论 220, 221, 223, 235
- 响度顺序原则 221, 223, 238
- 向 心 结 构 458, 470, 514, 515, 550
- 选择规则 99, 168, 286, 287, 292, 296, 300, 332
- 选择限制 87, 291, 300, 340, 342, 516
- 延展原则 226 — 228, 231 — 237, 242, 243
- 移位 50, 109, 125, 133, 139, 302, 306, 347, 391, 427, 439, 520, 542
- 音 段 145, 146, 154, 169, 173, 181, 189, 192, 194, 196 — 198, 216, 218 — 221, 223, 225, 227, 229, 230, 232 — 236, 238 — 240, 242, 243, 245, 247, 249, 250, 252 — 255, 260, 264, 277
- 音 节 52, 54, 58, 59, 75, 141, 149 — 152, 171, 172, 184 — 186, 189, 192, 200, 201, 204, 206, 207, 209, 210, 212, 213, 215 — 221, 223 — 225, 227 — 229, 231 — 234, 236 — 240, 242, 244, 249, 252, 267, 272, 277, 307, 424, 492, 507 — 509
- 音节结构 240, 243
- 音强音节观念 223, 235

- 音位 145, 147, 149, 151, 152, 154, 155, 157, 158, 160—163, 166, 167, 169, 173, 175, 176, 178, 179, 182, 184, 186—188, 192, 198, 202, 206, 216, 238, 240, 243, 245, 246, 248, 249, 251—254, 256, 260, 261, 267, 268, 272, 277, 281, 305
- 音子 155, 156, 158, 175, 188, 192, 193, 196—198, 217, 246, 247, 249, 254, 256, 258, 260, 261, 266—268, 272, 276
- 有限状态语法 81—83, 85—88, 92—94, 96, 98, 100, 104, 105, 110, 113, 115, 120
- 有限状态自动机 81, 83, 85, 86, 89—92
- 语法功能 138, 263, 319, 360, 361, 365, 366, 460
- 语符 23, 47, 50, 54, 55, 57—59, 61, 63, 64, 67, 73, 76, 77, 80, 84, 87, 88, 92, 94, 100—102, 105, 110, 111, 113, 116, 117, 120, 144, 285, 307, 317, 357, 358, 360, 368, 370, 373, 376, 377, 387, 391—394, 396, 417, 420, 422, 432, 433, 435, 442, 447—449, 496, 500, 501, 509, 519
- 语符库 50, 54, 56, 58, 60, 61, 64, 67, 68, 73, 76, 141, 165, 166, 210—212, 307, 353, 357, 358, 360, 371, 375, 395, 432, 435
- 语素 1, 5, 7, 9—13, 15, 16, 32, 35—39, 42, 44—50, 52—55, 57—67, 73, 78, 79, 81, 118, 121—124, 128, 141, 145—155, 158, 161, 163—171, 173—186, 188—190, 192, 193, 195, 196, 198—202, 204—207, 209—214, 216, 240—250, 252, 254, 256, 268, 272, 299, 307, 561
- 语素本音 146, 205
- 语素基本音形 146, 155, 158, 162, 171, 182, 196, 202
- 语言能力 290, 567
- 语义表达式 284, 302, 399, 401, 405—408, 410, 411, 414, 419, 420, 432, 434—436, 446, 454
- 语义格 128, 308, 319—321,

- 324, 326 — 331, 333 — 336,
338 — 345, 347 — 351, 353 —
358, 360, 363, 367, 371,
373, 377, 386 — 389, 391,
393 — 395, 403, 408, 410,
413, 432, 434, 436, 454,
457, 458, 521, 555, 556
- 语义功能 319, 320, 331,
335, 342, 345, 435, 556, 570
- 语义特征 54, 55, 288, 290,
292, 294, 325, 334, 336,
341, 342, 345, 349, 359,
373, 375, 387, 388, 390
- 域内论元 363
- 域外论元 360, 363
- 元语言 104, 409, 410, 413,
415 — 417, 419, 422, 423,
425
- 原则与参数 531
- 约束理路 520, 521
- 约束理论 302, 387, 521,
524, 553
- 暂拟语素 39, 43
- 照应成分 521
- 支配范畴 522 — 526, 528,
530, 532
- 支配理论 38, 41, 532, 547,
549, 550
- 直接宾语 554
- 直接成分 12, 15, 26 — 28,
31, 33, 36, 38, 48, 50, 78,
80, 81, 84, 85, 88, 91 — 95,
101, 103, 105, 108, 111,
123, 137, 139 — 143, 239,
272, 285, 291, 294, 299,
312, 319, 331, 364, 405,
458, 461, 465, 466, 468,
469, 471, 476, 478, 480,
489, 512, 514, 519, 546,
548, 550, 553, 555, 557, 570
- 直接成分控制 534, 538
- 直接成分同标法 465
- 直接关联原则 285, 286,
287, 288, 289, 290, 291,
292, 304
- 指称 56, 58, 303, 347, 384,
397, 421, 422, 441, 474,
475, 502, 521, 522, 524,
526, 528 — 532, 536 — 538
- 中动句 340
- 中和 160, 194, 357, 512
- 主语 1, 103, 315, 319, 325,
327, 330 — 332, 334 — 337,
339 — 342, 346, 347, 424,
426 — 428, 430, 437, 441,
454, 520, 526, 529, 530,
549, 554, 556
- 转换 2, 3, 5, 28, 30 — 32, 36,

- 38, 40, 46, 47, 49, 66, 68 —
71, 73, 80, 91, 92, 106 —
109, 122, 123, 125 — 128,
130 — 133, 135, 137, 139,
140, 143, 144, 170, 215,
247, 282, 284, 300, 302,
305, 307, 313, 315, 320,
322, 325, 330, 333, 343,
347, 355, 367, 374, 376,
379, 381, 382, 384, 391,
404, 406, 407, 409, 411 —
413, 426 — 428, 434 — 438,
440, 443, 445, 449 — 451,
453 — 456, 480, 482, 502,
510, 532, 546, 556, 560,
561, 563, 566
- 转换假说 47
- 转换生成语法 2, 38, 47, 80,
102, 123, 413, 446, 449
- 姊妹关系 468, 549, 551
- 自立音系学 161, 162
- 自由语义格 351, 352
- 自主音段音系学 196
- 总体分布 17, 18
- 组合关系 1, 5, 6, 16, 36, 51,
53, 55, 57, 80, 85, 86, 99,
100, 104, 173, 282, 292,
319, 350, 515
- 组合基式 15, 22
- 最大投射 462, 466 — 468,
475, 511, 538 — 540, 544,
545
- 最简方案 303, 309
- 作格 335, 336

参考书目

本书在讨论一些关键性理论问题时,为了不引起理解上的争议,尽量使用原文。一般情况下直接使用汉译文。凡是外文书目,译本和原文都列出,读者可根据人名译名对照表查阅。译文主要根据译本,若和译本有出入,都是根据原文的意思更改的,并加以说明。外文文献没有译本的,一般情况下直接给出译文。请读者能尽量阅读原文。

作者后面的第1个年代是作者最早发表论著的年代,也是书中人名后面的括号中的年代。出版社后面的年代是本书引用的版本年代,两者相同时后一个年代省略。方括号中的内容是该文献和本书内容有联系的主要观点或内容,涉及范围较多的不再注明,请读者从本书相关论述中查找。

Anderson J. M. , 1977, *On Case Grammar* , Groom Helm, London.

Aoun J. , Sportiche D. , 1981, *On the formal theory of government* , The Linguistic Review.

Austin J. L. , 1962, *How to Do Things with Words* , Clarendon Press, Oxford.

Bao Zhiming, 1990, *On the Nature of Tone* , Ph. D. dissertation. MIT. Cambridge, Mass.

Bell A. , Hooper J. B. (eds), 1978, *Syllables and Segments* , North-Holland, Amsterdam.

Bloch B. , 1941, *Phonemic overlapping* , in Joos Martin (1957).

- Bloomfield L. ,1926, *A set of postulates for the Science of language*, Language Vol. 2.
- Bloomfield L. ,1933, *Language*, Henry Holt, New York.
- Chen M. Y. , 1979, *Metrical structure: Evidence from Chinese poetry*, Linguistic Inquiry 10.
- Cheng C. C. ,1973, *A Synchronic Phonology of Mandarin Chinese*, Mouton, The Hague. [用生成音系学描写汉语]
- Chomsky N, Halle M. , Lukoff F. ,1956, *On accent and juncture in English*, For Roman Jakobson, Mouton, The Hague, 65—80.
- Chomsky N. , 1955, *The Logical Structure of Linguistic Theory*, Plenum, New York, 1975. [形式语法, 详细讨论了转换规则代数和转换规则语法]
- Chomsky N. , 1956, *Three models for the description of language*, I. R. E. Transactions on Information Theory, Vol. IT-2, Proceedings of the Symposium on Information Theory, Sept, 1956. [形式文法]
- Chomsky N. , 1957, *Syntactic Structures*, Mouton, the Hague.
- Chomsky N. , 1964, *Current Issues in Linguistic Theory*, Mouton, the Hague. [讨论了描写的充分性和解释的充分性]
- Chomsky N. , 1965, *Aspects of the Theory of Syntax*, MIT Press, Cambridge, Mass.
- Chomsky N. , 1970, *Remarks on nominalization*, in R. Jacobs & P. Rosenbaum, eds. , *Readings in English Transformational Grammar*, 184—221. Ginn, Waltham, Mass.
- Chomsky N. , 1972, *Deep structure, surface structure and*

- semantic interpretation*, In N. Chomsky, *Studies on Semantics in Generative Grammar*, 11—61 Mouton, the Hague. [生成语法扩展标准理论]
- Chomsky N. , 1973, *Conditions on transformations*, In S. Anderson & P. Kiparsky, eds. , *A Festschrift for Morris Halle*, 232 - 286. [讨论转换条件]
- Chomsky N. , 1981, *Lectures on Government and Binding*, Foris Publications, Holland, 1982.
- Chomsky N. , Miller G. A. , 1958, *Finite state languages*, *Information and Control*, Vol. 1:2. [形式文法]
- Chomsky N. , Halle M. , 1968, *The Sound Pattern of English*, Harper & Row, Publishers, New York. [生成音系学标准理论]
- Chomsky N. , 1956b, *Three models for the description of language*, *PGIT*, 2:3, 113—124.
- Chomsky N. , 1986a, *Knowledge of Language: It's Nature, Origin and Use*, Praeger, New York.
- Chomsky N. , 1986b, *Barriers*, MIT Press, Cambridge, Mass.
- Chomsky N. , 1995, *The minimalist program*, MIT Press, Cambridge, Mass.
- Chomsky N. , Lasnik H. , 1977, *Filters and control*, *Linguistic Inquiry* 8, 425—504.
- Church A. , 1936, *An unsolvable problem of elementary number theory*, *Amer. J. Math.* , 58, 345—363. [形式文法]
- Clark J. , Yallop C. , 1990, *An Introduction to Phonetics and Phonology*, 外语教学与研究出版社, 2000.
- Clements N. , 1990, *The role of the sonority cycle in core*

- syllabification*, In J. Kingston & M. Beckman (ed.).
Papers in Laboratory Phonology 1: Between the Grammar and Physics of Speech. Cambridge: Cambridge University Press. [音节响度理论]
- Dinssen D. (ed.), 1979, *Current approaches to phonological theory*, Bloomington and London: Indiana University Press.
- Dowty David R., 1991, *Thematic Proto-Roles and Argument Selection*, *Language* 67: 547—619. [论元角色理论]
- Duanmu San, 1990, *A Formal Study of Syllable, Tone, Stress and Domain in Chinese*, *MIT working papers in Linguistics*, Cambridge, Mass.
- Duanmu San, 1995, *Metrical and phonology of compounds in two Chinese dialects*, *Language* 71. 2: 225—259.
- Durand J., 1990, *Generative and Non-linear Phonology*, Longman, London. [给出了响度阶]
- Feng Shengli, 1990a, *The passive construction in Chinese*, ms., University of Pennsylvania. [空语类]
- Feng Shengli, 1990b, *Subject in Chinese and the theory of case-assignment*, *PENN Reveiw of Linguistics* Vol. 14. [空语类分析]
- Ferguson C., 1962, *Review of Halle, The sound pattern of Russian*, *Language* 38. 284—297.
- Fillmore C., 1963, *The position of embedding transformations in a grammar*, *Word*, 19. PP. 208—231.
- Fillmore C., 1966, *A proposal concerning English prepositions*, *Monograph Series on Languages and Linguistics* 19. 19—34.
- Fillmore C., 1968, *The case for case*, In E. Bach and E.

- Harms, eds. , *Universals in Linguistic Theory*, 1—90.
[提出格语法理论]
- Fodor J. A. , Katz J. J. (eds.), 1964, *The Structure of Language: Readings in the Philosophy of Language.* , Englewood Cliffs, N. J. : Prentice-Hall.
- Fraser B. , 1963, *The position of conjoining transformations in a grammar*, Mimeographed. Mitre Corporation, Bedford, Mass.
- Gazdar G. , 1982, *Phrase structure grammar*, In Jacobson, P. & G. Pullum, eds. , *The Nature of Syntactic Representation*, Dordrecht: Reidel. 131—186.
- Gazdar G. etc. , 1985, *Generalized Phrase Structure Grammar*, Harvard University Press. Cambridge, Mass. [全面论述广义短语结构语法]
- Goldberg A. E. , 1995, *Construction: A Construction Grammar Approach to Argument Structure*, Chicago University Press, Chicago.
- Goldsmith J. , 1976, *Autosegmental Phonology*, PhD dissertation, MIT, Cambridge, Mass. [自主音段音系理论, 多线性分析]
- Goldsmith J. , 1979, *The aims of autosegmental phonology*, In Dinnsen 1979: 202—222. [自主音段音系理论, 多线性分析]
- Gruber J. S. , 1965, *Studies in Lexical Relations*, Ph. D dissertation, MIT, Cambridge, Mass. [转换生成语法中早期语义格的研究]
- Gruber J. S. , 1976, *Lexical Structures in Syntax and Semantics*, North Holland, Amsterdam.
- Halle M. , Vergnaud J. R. , 1978, *Metrical structures*

- in phonology*, (A fragment of a draft). ms. [节律音系理论]
- Halle M. , Vergnaud J. R. , 1980, *Three dimensional phonology*, *Journal of Linguistic Research* 1. 83—100.
- Halle M. , 1959, *The Sound Pattern of Russian*, Mouton, The Hague.
- Halliday M. , 1985, *An introduction to functional grammar*, Edward Arnold, London. [功能语法研究]
- Harris Z. S. , 1942, *Morpheme alternants in linguistic analysis*, Joos, 1957. [语素的线性切分和归并]
- Harris Z. S. , 1945, *Discontinuous morphemes*, *Language* 21. [不连续语素]
- Harris Z. S. , 1946, *From morpheme to utterance*, *Language* 22, 166—183. 译文载《语言学资料》1963. 6. [通过语素分布描写句法]
- Harris Z. S. , 1951, *Methods in Structural Linguistics*, University of Chicago Press, Chicago. [全面讨论分布分析]
- Harris Z. S. , 1955, *From phoneme to morpheme*, *Language* 31, No. 2, 190—222. [从音位出发描写语素]
- Harris Z. S. , 1964, *String Analysis of Sentence Structure*, Mouton, The Hague. [提出了分析话语的线性分析法]
- Harris Z. S. , 1965, *Transformational theory*, *Language* 41, No. 3, Part 1. [把 transformation 作为分析方法进行理论分析]
- Harris Z. S. , 1952, *Discourse analysis*, *Language* 28. [首次提出了 transformation 的分析方法]
- Harris Z. S. , 1957, *Co-occurrence and transformation in*

- linguistic structure*, *Language* 33. [可观察结构的转换分析]
- Hartman L. M., 1944, *The segmental phonemes of the Peking Dialect*, *Language* 20. [用 IA 模式或纯线性模式描写北京话]
- Hayes B., 1984, *The Phonology of Rhythm in English*, *Linguistic Inquiry*, Vol.
- Hockett C. F., 1942, *A system of descriptive phonology*, *Language* 18.
- Hockett C. F., 1947, *Peking phonology*, *Journal of American Oriental Society* 67. [用 IA 模式描写北京话]
- Hockett C. F., 1954, *Two Models of Grammatical Description*, *Word* 10. [概括出结构语言学的 IA 和 IP 两种描写模式]
- Hockett C. F., 1958, *A Course in Modern Linguistics*, The Macmillan Company, New York, 1968. [结构语言学集大成著作]
- Hockett C. F., 1961, *Linguistic elements and their relations*, *Language* 37.
- Hooper J. B., 1976, *An Introduction to Natural Generative Phonology*, Academic Press, New York. [不主张使用不可观察的底层音系结构]
- Hopper J. B., 1979, *Substantive principles in Natural Generative Phonology*, In Dinnsen 1979: 106-125. [不主张使用不可观察的底层音系结构或单位]
- Huang C. T. James, 1982, *Logical Relations in Chinese and the Theory of Grammar*, PhD dissertation, MIT, Cambridge, Mass. [汉语生成语法]
- Huang C. T. James, 1987, *Remarks on empty categories in*

- Chinese*, *Linguistic Inquiry*, 18. [汉语空语类研究]
- Huang Yunhua, 1984, *Reflexives in Chinese*, *Studies in English Literature and Linguistics* 10.
- Jackendoff R. S., 1969a, *Some rules of semantic interpretation for English*, PhD dissertation, MIT, Cambridge, Mass.
- Jackendoff R. S., 1969b, *An Interpretive Theory of Negation*, *Foundations of Language* V.
- Jackendoff R. S., 1977a, *X-bar Syntax: A Study of Phrase Structure*, MIT Press, Cambridge, Mass.
- Jakobson R., C. Fant and M. Halle, 1952, *Preliminaries to Speech Analysis*, The MIT Press, Cambridge, Mass.
- Jespersen Otto, 1913, *Lehrbuch der Phonetik* (语音学), B. G. Teubner, Leipzig. [早期的音节响度理论]
- Jespersen Otto, 1922, *Language: its nature, development and origin*, Allen and Unwin, London.
- Joos M. ed., 1957, *Readings in Linguistics*, American Council of Learned Societies, Washington. [美国描写语言学论文选, 1925—1957]
- Katz J. J., Fodor J. A., 1963, *The structure of a semantic theory*, *Language* 36, Reprinted in Fodor and Katz (1964). [提出: 语言描写—语法学=语义学]
- Katz J. J., Postal P. M., 1964, *An Integrated Theory of Linguistic Descriptions*, MIT Press, Cambridge, Mass. [语法学中应包括语义部分, 并在深层结构解释语义]
- Keyser S. J., Roeper T., 1984, *On the middle and ergative constructions in English*, *Linguistic Inquiry*, 15.

- pp. 381—416. [中动和作格现象研究]
- Kiparsky P. , 1982, *From cyclic phonology to lexical phonology*, In H. van der Hulst and N. Smith (eds). *The Structure of Phonological Representations Part 2*. Foris, Dordrecht. [词汇音系学理论]
- Klima E. S. , 1964, *Negation in English*, in Fodor and Katz (1964).
- Ladefoged, 1967, *Three Areas of Experimental Phonetics*, Oxford University Press, London.
- Lakoff G. , 1970. , *Global Rules*, *Language* 46. [派生的普遍性问题]
- Lakoff G. , 1971. , *On Generative Semantics*, in *Semantics*, edited by Steinberg, D. D. and Jakobovits, A. , Cambridge University Press.
- Lakoff G. , Ross J. R. , 1967, *Is deep structure necessary*, In J. McCawley, ed. , *Syntax & Semantics* 7. [取消深层结构]
- Langacker R. W. , 1987, *Foundations of Cognitive Grammar*, Vol. I. Stanford University Press, Stanford.
- Langacker R. W. , 1969, *On pronominalization and the chain of command*, In D. Reibel and S. Schane, ed., *Modern studies in English*. Englewood Cliffs, N. J. : Prentice-Hall.
- Larson R. , 1988, *On the Double Object Construction*, *Linguistic Inquiry* 19.
- Larson R. , 1990, *Double Object Revisited: A Reply to Jackendoff*, *Linguistic Inquiry* 21.
- Lasnik H. , 1976, *Remarks on coreference*, *Linguistic Analysis* 2, 1—22, Reprinted in Lasnik 1989.

- Lasnik H. ,1989, *Essays on anaphora* ,Dordrecht: Reidel.
- Lees R. B. ,1960, *The Grammar of English Nominalizations* ,Mouton, The Hague.
- Levin B. , Hovav M. R. ,2005, *Argument Realization* , Cambridge University Press, Cambridge.
- Li C. N. , Thompson S. A. ,1976, *Subject and topic: a new topology of language* ,In Li ed. , Subject and Topic. Academic Press, New York.
- Martinet A. ,1960, *E'lémennts de linguistique générale* ,Armand colin 1973, Paris.
- McCawley J. ,1968, *Lexical insertion in Transformational Grammar without deep structure* ,Papers from the 4th regional meeting of the Chicago Linguistic Society.
- Nespor M. , Vogel I. ,1986, *Prosodic Phonology* ,Foris Publications, Holand.
- Perlmutter D. , Postal P. ,1977, *Toward a universal characterization of passive* , Proceedings of the Third Annual Meeting of the Berkeley Linguistics Society.
- Postal P. M. ,1962, *On the limitations of context-free phrase-structure description* ,Quarterly Progress Report, No. 64, Research Laboratory of Electronics, MIT, Cambridge, Mass.
- Postal P. M. ,1964, *Limitation of phrase structure grammars* ,In J. A. Fodor & J. J. Katz, eds. , The Structure of Language.
- Radford A. ,1981, *Transformational Syntax-A student's guide to Chomsky's extended standard theory* , Cambridge University Press, Cambridge.
- Radford A. ,1988, *Transformational Crammar: A First*

- Course*, Cambridge University Press, Cambridge.
- Reinhart T. , 1976, *The Syntactic Domain of Anaphora* ,
Ph. D dissertation, MIT, Cambridge, Mass.
- Rosetti A. , 1961, *Sur le probleme de la syllabe*, *Phonetica* 7.
- Ross J. , 1967, *Constraints on Variables in Syntax* , Ph. D.
dissertation, MIT, Cambridge, Mass.
- Ryle G. , 1953, *Ordinary language* , *Philosophy Review* ,
LXII(1953).
- Sapir E. , 1921, *Language*, Rupert Hart-Davis.
- Selkirk E. , 1978, *On Prosodic Structure and its Relation to
Syntax Structure* , In *Nordic Prosody II* , Fretheim, T.
ed. , Trondheim.
- Selkirk E. , 1984. *Phonology and Syntax* , MIT Press,
Cambridge, Mass.
- Selkirk E. , 1984a, *On the major class features and syllable
theory*, In M. Liberman & R. T. Oehrle (eds.). *Language Sound Structure*, MIT Press, Cambridge, Mass.
- Shannon C. E. , Weaver W. , 1949, *The Mathematical
Theory of Communication* , Urbana.
- Stetson R. H. , 1924, *Motor phonetics* , North-Holland,
Amsterdam, 1951.
- Stowell T. , 1981, *Origins of phrase structure* , PhD Disser-
tation, MIT, Cambridge, Mass.
- Swadesh, 1934, *The phonemic principle* , *Language* vol. 10.
[提出声调音位]
- Sweet H. , 1875—1876, *Word, logic, grammar*, *Transac-
tions of the Philological Society*, 译文载《中国语文》
1961. 9. [词能构成独立的句子]
- Tang C.-C. J. , 1989, *Chinese Reflexives* , *Natural Lan-*

- guage and Linguistic Theory 7.
- Tesnière L. ,1934, *Comment construire une syntaxe*, Bulletin de la Faculté des Lettres de Strasbourg. [首次阐释了从属关系语法的基本论点]
- Tesnière L. ,1959, *E'léments de Syntaxe Structurale*, Paris. [提出了价的概念]
- Trubetskoy N. S. ,1939, *Principles of Phonology*, University of California Press, 1969. [全面讨论了音位理论]
- Trubetzkoy N. S. ,1969, *Principles of Phonology*, University of California Press. [全面讨论了音位理论]
- Van Valin R. , Lapolla R. , 1979, *Syntax: Structure, Meaning and Function*, Cambridge University Press, Cambridge.
- Vennemann T. , 1974, *Phonological concreteness in natural generative phonology*, In Shuy and Bailey 1974:202—19.
- Wang William S-Y. ,1964, *Some syntactic rules for Mandarin*, Proceedings of the Ninth International Congress of Linguists, Mouton. [根据分布给动词分类]
- Wang William S-Y. ,1967, *Phonological features of tone*, International Journal of American Linguistics (译文载石锋《语音学探微》). [调位的独立性]
- Wells R. S. ,1947, *De Saussure's system of linguistics*, Word 3.
- Wells R. S. ,1947, *Immediate constituents*, Language 23, also in Joos 1957.
- William E. ,1981, *Argument structure and morphology*, The Linguistic Review 1.
- Xu L. J. ,1986, *Free empty categories*, Linguistic Inquiry 17. [空语类]
- Xu L. J. ,1993, *The Long-Distance Binding of Ziji*, Jour-

nal of Chinese Linguistics.

Xu L. J. and D. T. Langenden, 1985, *Topic structures in Chinese*, Language 61. [根据汉语施事提出对空语类的修正]

Zec D, 1988, *Sonority Constraints on Prosodic Structure*, PhD Dissertation, Stanford University.

Zhang Hongming, 1992, *Topics in Chinese Phrasal Tonology*, PhD dissertation. UC, San Diego. [对韵律音系学的介绍研究]

包智明, 1997, 《从晋语分音词看介音的不对称性》, 《中国语言学报》第 1 辑, 北京语言文化大学出版社。

北京大学中文系现代汉语教研室, 1993, 《现代汉语》, 商务印书馆。

布龙菲尔德, 1933, 《语言论》, 袁家骅等译, 商务印书馆 1980。

曹剑芬, 1995. 4, 《连读变调与轻重对立》, 《中国语文》。

查理斯、李纳, 1984. 2, 《主语与主题: 一种新的语言类型学》, 《国外语言学》。

岑麒祥, 1955. 10, 《讨论主语宾语问题的几个原则》, 《语文学习》。

陈保亚, 1988, 《语言演变的结构基础》, 北京大学中文系硕士学位论文, 部分内容载《缀玉集》第一集, 北京大学出版社 1990。

陈保亚, 1992, 《语言文化论》, 云南大学出版社。

陈保亚, 1999, 《20 世纪中国语言学方法论》, 山东教育出版社。

陈保亚, 2005. 1, 《再论平行周遍原则和不规则字组的判定》, 《汉语学习》。

陈保亚, 2006. 2, 《论平行周遍原则与规则语素组的判定》, 《中国语文》。

- 陈立民,1995,《论汉语动词配价分类的原则》,北京大学中文系硕士学位论文。[设法先建立汉语的格系统,再通过这个系统来确定配价原则]
- 陈平,1987.5,《汉语零形回指的话语分析》,《中国语文》。[用参与成分(participant role)和附带成分(circumstantial role)确定价]
- 陈平,1991,《现代语言学研究——理论、方法与事实》,重庆出版社。
- 程工,1994.1,《生成语法对“自己”一词的研究》,《国外语言学》。
- 程工,1999.2,《汉语“自己”一词的性质》,《当代语言学》。
- 崔延燕,2008,《从事件结构看现代汉语述宾式复合动词的论元实现》,北京大学中文系硕士论文。
- 戴浩一,1988.1,《时间顺序和汉语的语序》,《国外语言学》,黄河译。[汉语语序和事件先后发生的一致性]
- 戴庆厦、罗仁地、汪锋,2008,《到田野去——语言学田野调查的方法与时间》,民族出版社。
- 丹尼尔·琼斯,1980.2,《音位的历史和涵义》,《国外语言学》。
- 丁崇明,1992.1,《大理方言中与动词给相关的句式》,《中国语文》。
- 丁声树等,1952.7-1953.11,《现代汉语语法讲话》,商务印书馆1961。[用结构主义方法描写汉语]
- 董秀芳,2004,《汉语的词库和词法》,北京大学出版社。
- 端木三,1997,《从汉语的重音谈语言的共性与特性》,《中国语言学论丛》第1辑。
- 范继淹,1958.5,《形名组合间“的”字的语法作用》,《中国语文》。[分布分析]
- 范继淹,1963.2,《动词和趋向性后置成分的结构分析》,《中国语文》。[提出了层次分析的并立扩展法]
- 范继淹,1983,《汉语语法结构的层次分析问题》,《语法研究和探

- 索》(一),北京大学出版社。[讨论了并立扩展方法]
- 范晓,1991,《动词的价分类》,《语法研究和探索》(五),语文出版社。[用强制性共现标准确定价。借助介词给动词定价]
- 方德义,1986.3,《法国现代语言学理论研究概况》,《国外语言学》。
- 方经民,1989.1—2,《哈里斯的变换理论》,《语言学通讯》。
- 方立,1993,《美国理论语言学研究》,北京语言学院出版社。
- 方梅,1993.1,《宾语与动量词语的次序问题》,《中国语文》。
- 方希,2002,《黏合式多重定名结构的顺序》,《语言学论丛》第25辑。
- 菲尔墨,1968,《格辨》,胡明扬译,《语言学译丛》第2辑,中国社会科学出版社1980。[提出格语法理论]
- 冯胜利,1997,《管约理论与汉语的被动句》,《中国语言学论丛》第1辑,北京语言文化大学出版社。[转换或移位]
- 冯志伟,1979.1,《形式语言理论》,《计算机科学》,34—57。
- 冯志伟,1983.1,《特思尼耶尔的从属关系语法》,《国外语言学》。
- 冯志伟,1987,《现代语言学流派》,陕西人民出版社。
- 符淮青,1996,《词义的分析和描写》,语文出版社。[提出词义成分——模式分析法]
- 弗雷格,1988,《论涵义与指称》,肖阳译,《语言哲学名著选辑》,三联书店。
- 傅懋勳,1956.5,《北京话的音位和拼音字母》,《中国语文》。[声调是元音音位的一个构成部分]
- 高名凯,1948,《汉语语法论》,上海开明书店出版。[提出了汉语中的一些独立的范畴]
- 郭锐,1995,《述结式的配价结构与成分的整合》,沈阳、郑定欧主编《现代汉语配价语法研究》,北京大学出版社。[提出了述结式配价的公式]
- 郭锐,2002,《现代汉语词类研究》,商务印书馆。

海里斯,1946,《从语素到话语》,《语言学资料》1963.6。[通过语素分布描写句法]

海里斯,1963.6,《语言分析中的语素交替形式》,《语言学资料》。[语素的线性切分和归并]

贺巍,1979.2,《获嘉方言的连续变调》,《方言》。

贺阳,2008.4,《现代汉语欧化语法现象研究》,《世界汉语教学》。

侯精一,1985.2,《晋东南地区的子变韵母》,《中国语文》。

胡炳忠,1985.1,《三声三字组的变调规律》,《语言教学与研究》。

胡建华,1998.3,《汉语长距离反身代词化的句法研究》,《当代语言学》。

胡明扬,1991.2,《句法语义范畴的若干理论问题》,《语言研究》。

胡裕树,1982.4,《试论汉语句首的名词性成分》,《语言教学与研究》。[语用分析]

胡壮麟、朱永生、张德录,1994,《系统功能语法》,湖南教育出版社。

黄国营,1982.1,《“的”字的句法、语义功能》,《语言研究》。[语义特征分析,分布分析]

黄正德,1983,《汉语生成语法》,黑龙江大学出版社。

霍凯特,1950,《北京话形态音素学》,《国外语言学》1980.5—6。

霍凯特,1954,《语法描写的两种模式》,《语言学资料》1963.6。[概括出结构语言学的 IA 和 IP 两种描写模式]

霍凯特,1958,《现代语言学教程》,索振羽、叶蜚声译,北京大学出版社 1986。[结构语言学的集大成著作]

霍凯特,1961,《语言的各种单位及其关系》,《语言学资料》1964.1。

霍凯特,1963.6,《语素分析的一些问题》,《语言学资料》。

霍普克罗夫特、厄尔曼,1969,《形式语言及其与自动机的关系》,莫绍揆等译,科学出版社 1979。[形式文法]

卡尔纳普,1934,《哲学和逻辑句法》,上海人民出版社 1962。

李临定,1983.2,《宾语使用情况考察》,《语文研究》。[对汉语的

语义格进行了比较全面的考察]

李小凡,1990,《苏州话的字调转移及其成因》,《缀玉集》,北京大学出版社。[结构与变异。不同层面的干扰]

林焘、王理嘉,1992,《语音学教程》,北京大学出版社。

林焘、沈炯,1995.3,《北京话儿化韵的语音分歧》,《中国语文》。
[变异分析]

刘丹青,1987.3,《形名同现及形容词的向》,《南京师范大学学报》。

刘宁生,1995.2,《汉语偏正结构的认知基础及其在语序类型学上的意义》,《中国语文》。[用人类认知的普遍语义或逻辑概念来描写和解释汉语语法规则]

刘现强,2002.10,《现代汉语节奏研究》,北京语言语言文化大学出版社。

刘月华,1983,《状语的分类和多项状语的顺序》,《语法研究和探索》(一),北京大学出版社。[状语的语义指向]

龙果夫,1958,《现代汉语语法研究》,郑祖庆译,科学出版社。
[根据词汇·语法范畴划分词类]

鲁川、林杏光,1989.5,《现代汉语语法的格关系》,《汉语学习》。
[给语义格分层次]

陆丙甫,1979.4,《读〈的字结构和判断句〉》,《中国语文》。[提出组合关系的三个层面]

陆丙甫,1989.3,《结果、节奏、松紧、轻重在汉语中的相互作用》,《汉语学习》。

陆俭明,1980.1,《汉语口语句法里的易位现象》,《中国语文》。
[区分了语义结构关系和语法结构关系]

陆俭明,1990.3,《变换分析在汉语语法研究中的运用》,《湖北大学学报》。

陆俭明,1997,《关于语义指向分析》,《中国语言学论丛》第1辑。
[总结了语义指向分析]

- 陆志韦,1937,《北京话单音词词汇》,人民出版社 1951,又载《陆志韦语言学著作集》(三)。[提出同形替代法分离词]
- 陆志韦,1957,《汉语构词法》,《陆志韦语言学著作选》(三),中华书局 1990。[比较深入全面地讨论了扩展法]
- 陆致极,1985.3—4,《关于非线性音位学》,《国外语言学》。
- 吕叔湘,1942,《中国文法要略》,商务印书馆 1982。[首次从造句的角度比较明确地提出了转换的概念,重视语义结构关系的分析]
- 吕叔湘,1946,《从主宾语的分别谈国语句子的分析》,吕叔湘《汉语语法论文集》,科学出版社 1955。[有用施受关系取消主宾关系的倾向]
- 吕叔湘,1954,9—10,《关于汉语词类的一些原则性问题》,《中国语文》。[讨论了用鉴定字划分词类的方法]
- 吕叔湘,1962.1,《说自由和粘着》,《中国语文》。
- 吕叔湘,1962.11,《关于语言单位的同一性等等》,《中国语文》。[分布分析]
- 吕叔湘,1979,《汉语语法分析问题》,商务印书馆。[讨论了汉语语法研究中的很多基本方法论问题]
- 吕叔湘(主编),1980,《现代汉语八百词》,商务印书馆。[分布分析]
- 罗曼玲,1998,《现代汉语同义词的词义分析和组合分析》,北京大学中文系硕士学位论文。
- 罗素,1905,《论指谓》,牟博译,《语言哲学名著选辑》,三联书店 1988。[摹状词理论]
- 马庆株,1983,《现代汉语的双宾语构造》,《语言学论丛》第 10 辑。[认为价是指最小的主谓结构中动词(不借助于介词)所能联系的名词成分的数目]
- 马庆株,1988,《自主动词和非自主动词》,《中国语言学报》第 3 期 [语义特征分析]

- 马庆株,1992,《汉语动词和动词性结构》,语言学院出版社。[语义特征分析]
- 帕默尔,1936,《语言学概论》,商务印书馆 1986。
- 钱军,1998. 12,《结构功能语言学——布拉格学派》,吉林教育出版社。
- 钱乃荣,1988. 1,《论普通话语音的音位和区别性特征》,《汉语学习》。[按声韵调归纳音位,得声位、韵位和调位]
- 钱曾怡,1963. 1,《济南话的变调和轻声》,《济南话的变调和轻声》。
- 钱曾怡,1995,《论儿化》,《中国语言学报》第 5 期。
- 乔姆斯基,1957,《句法结构》,邢公畹等译,中国社会科学出版社 1979。
- 乔姆斯基,1965,《句法理论的若干问题》,黄长著等译,中国社会科学出版社 1986。
- 乔姆斯基,1972,《深层结构、表层结构和语义解释》,赵世开译,《语言学译丛》第 2 辑,中国社会科学出版社 1980。[生成语法扩展标准理论]
- 乔姆斯基,1981,《支配和约束论集》,周流溪、林书武、沈家煊等译,中国社会科学出版社 1993。
- 邱立坤,2004,《现代汉语偏正式动名短语研究》,北京大学中文系硕士论文。
- 荣晶,1988,《汉语省略、隐含和空语类的区分》,新疆大学中文系硕士学位论文,部分内容载《新疆大学学报》1989. 4。
- 荣晶,1997,《汉语语序的语义基础》,北京大学中文系博士论文。[语义范畴对语序的制约]
- 萨丕尔,1921,《语言论》,商务印书馆 1985。
- 邵敬敏,1995,《双音节 V+N 结构的配价分析》,沈阳、郑定欧《现代汉语配价语法研究》。
- 申小龙,1986. 1,《汉语动词的分类角度》,《语言教学与研究》。

- 沈家煊,1993.5,《语用否定考察》,《中国语文》。[语用分析]
- 沈家煊,1995.5,《有界与无界》,《中国语文》。[提取了有界和无界的语义范畴]
- 沈炯,1994.4,《北京话上声连读的调型组合和节奏形式》,《中国语文》。
- 沈开木,1984.6,《“不”字的否定范围和否定中心的探索》,《中国语文》。[“不”的语义指向]
- 沈阳,1994,《现代汉语空语类研究》,山东教育出版社。[变换分析,并提出句位概念]
- 沈阳、郑定欧,1995,《现代汉语配价语法研究》,北京大学出版社。
- 沈阳、何远建、顾阳,2001,《生成语法理论与汉语语法研究》,黑龙江教育出版社。
- 石安石,1978.4,《汉语词组基本类型的鉴别问题》,《天津师院学报》。[讨论根据推导式确定词组类型的方法]
- 石安石,1993.1,《论语素的结合能力与一用语素》,《语文研究》。
- 石定栩,2002,《乔姆斯基的形式句法》,北京语言文化大学出版社。
- 石锋,1990,《试论语音的层次》,石锋《语音学探微》,北京大学出版社。[语音实验的对象是音子]
- 石毓智,1995.3,《论汉语的大音节结构》,《中国语文》。
- 斯米尔尼兹基,1952,《词的分离性》,载陆志韦《汉语构词法》,科学出版社 1957。[讨论了提取词的剩余法]
- 宋作艳,2001,《照应语自己的约束》,当代语言学课程研究报告。
- 宋作艳,2005.1,《控制“一”变调的相关因素分析》,《汉语学习》。
- 索绪尔,1916,《普通语言学教程》,高名凯译,商务印书馆 1986。
- 汪平,1995,《汉语方言四呼比较》,《中国语言学报》第5期。
- 王辅世,1963.2,《北京话韵母的几个问题》,《中国语文》。[声调附属于韵母]
- 王洪君,1994.1,《汉语常用的两种语音构词法》,《语言研究》。

[生成音系学的方法。字本位观念]

王洪君, 1994. 2, 《从字和字组看词和短语》, 《中国语文》。[字本位观念]

王洪君, 1994a, 《什么是音系的基本单位——谈本音与变音》, 《现代语言学》, 语文出版社。[本音与变音分析]

王洪君, 1994b, 《生成音系学的形成和发展》, 石锋编《海外中国语言学研究》。

王洪君, 1996. 3, 《汉语语音词的韵律类型》, 《中国语文》。[生成音系学的方法。字本位观念]

王洪君, 1999, 《汉语非线性音系学》, 北京大学出版社。

王红旗, 1995, 《动词述补结构配价研究》, 沈阳、郑定欧主编《现代汉语配价语法研究》。[给出了动结式配价的计算公式]

王嘉龄, 1987. 2, 《词汇音系学》, 《国外语言学》。

王静、王洪君, 1995, 《动词的配价与被字句》, 沈阳、郑定欧《现代汉语配价语法研究》。[语义特征分析, 配价]

王理嘉, 1988, 《普通话音位研究中的几个问题》, 《语文研究》。

王理嘉, 1991, 《音系学基础》, 语文出版社。

王理嘉, 1995, 《儿化韵语素音位的讨论》, 《中国语言学报》第5期。

王力, 1943, 《中国现代语法》, 中华书局 1954。

王力, 1953. 9, 《词和伪语的界限问题》, 《中国语文》1953年9期 3—8页。

王玲玲, 1995, 《动词的必用论元与动词的向》, 沈阳、郑定欧《现代汉语配价语法研究》。

王士元, 1967, 《声调的音系特征》, 载石锋《语音学探微》北京大学出版社 1990。[声调特征独立于音段特征]

王志洁, 1999, 《北京话的音节与音系》, 《共性与个性》, 北京语言文化大学出版社。

威尔斯, 1947, 《直接成分》, 《语言学资料》1963. 6。[确定直接成

分的方法]

维特根斯坦,1922,《逻辑哲学论》,商务印书馆 1962。[提出了图式论,把意义限制在指称中]

维特根斯坦,1953,《哲学研究》,汤潮、范光棣译,三联书店 1992。[意义即用法]

吴为章,1982.5,《单向动词及其句型》,《中国语文》。

吴宗济,1985,《普通话三字组变调规律》,《中国语言学报》第2期,商务印书馆。

项梦冰,1997,《连城客家话语法研究》,语文出版社。

邢福义,1981.2,《评暂拟汉语教学语法系统》,《中国语文》。

邢欣,1995,《致使动词的配价》,沈阳、郑定欧《现代汉语配价语法研究》。

熊正辉,1984.2,《怎样求出两字组的连读变调规律》,《方言》。
[制约变调条件有今音的语音环境、古音来历、语法结构三个方面]

徐烈炯,1979.4,《两种新的音位学理论》,《语言学动态》。

徐烈炯,1984,《移位、空语类与领属条件》,《哈尔滨生成语法讨论会论文集》,黑龙江大学出版社。[空语类]

徐烈炯,1984.2,《管辖与约束理论》,《国外语言学》。

徐烈炯,1988,《生成语法理论》,上海外语教育出版社。

徐烈炯,1994.5,《与空语类有关的一些汉语语法现象》,《中国语文》。

徐通锵,1991.3,《语义句法刍议》,《语言教学与研究》。[提出字本位和语义句法的概念]

徐通锵,1994.2,《字和汉语的句法结构》,《世界汉语教学》。[提出字本位论]

徐通锵,1994.3,《字和汉语研究的方法论》,《世界汉语教学》。
[字本位论]

徐通锵,1997,《语言论》,东北师范大学出版社。[全面讨论字本

位和语义语法]

徐通锵,1997.1,《有定性范畴和语言的语法研究》,《语言研究》。

[语义范畴]

徐通锵、王洪君,1996,《改革开放以来的中国理论语言学》,载《中国语言学现状与展望》,外语教学与研究出版社。

徐通锵、叶蜚声,1979.3,《五四以来汉语语法研究评述》,《中国语文》。

徐云扬,1988,《自主音段音韵学理论与上海声调变调》,《当代语言学》。

许国璋,1983.1,《关于索绪尔的两本书》,《国外语言学》。

薛凤生,1982.2,《论音变与音位结构的关系》,《语言研究》。

薛凤生,1986,《北京音系解析》,北京语言学院出版社。[区别本音和变音]

雅各布森,1952,《语音分析初探》,《国外语言学》,1981.3—4。

杨成凯,1986.1—3,《Fillmore 的格语法理论》,《国外语言学》。

叶斯泊森,1924,《语法哲学》(The Philosophy of Grammar),语文出版社 1988。

叶文曦,1996,《汉语字组的语义结构》,北京大学中文系博士学位论文。[讨论了语义构词理论]

叶向阳,1997,《把字句的致使性解释》,北京大学中文系硕士学位论文。[论证了高层次语义关系和低层次语义关系,认为动词的及物关系是一种低层次语义关系]

游汝杰、钱乃荣、高钰夏,1980.5,《论普通话的音位系统》,《中国语文》。[按声韵调归纳音位,得声位、韵位和调位]

袁毓林,1989.1,《论变换分析方法》,《汉语学习》。

袁毓林,1992.3,《现代汉语名词的配价研究》,《中国社会科学》。

袁毓林,1993,《准双向动词研究》,《现代汉语祈使句研究》,北京大学出版社。

袁毓林,2002.3,《论元角色的层级关系和语义特征》,《世界汉语

教学》。

詹卫东,2000,《面向中文信息处理的现代汉语短语结构规则研究》,清华大学出版社。

张伯江、方梅,1994.2,《汉语口语的主位结构》,《北京大学学报》。〔语用分析〕

张伯江、方梅,1995,《北京口语易位现象的话语分析》,《语法研究和探索》(七),商务印书馆。〔语用分析〕

张伯江、方梅,1996,《汉语功能语法研究》,江西教育出版社。〔语用分析〕

张国宪,1993.2,《谈隐含》,《中国语文》。〔配价研究〕

张国宪,1995,《论双价形容词》,沈阳、郑定欧《现代汉语配价语法研究》。〔讨论了形容词的价〕

张和友,2002.10,《从焦点理论看汉语分裂式判断句的生成》,北京大学中文系博士生资格考试论文。

张惠英,1979.4,《崇明方言的连读变调》,《方言》。〔变调和语法条件有关〕

张力军,1990.3,《论 NP1 + A + VP + NP2 格式中 A 的语义指向》,《烟台大学学报》。〔语义指向〕

张敏,1989,《认知语言学与汉语名词短语》,中国社会科学出版社。

张志公等,1956,《暂拟汉语教学语法系统》,人民教育出版社。

赵元任,1934,《音位标音法的多能性》,《中国现代语言学的开拓和发展》,清华大学出版社 1992。〔讨论了音位归纳的相对性和严式记音的意义〕

赵元任,1948, Mandarin Primer, An Intensive Course in Spoken Chinese (《北京口语语法》), Harvard Univ. Press。〔最早用结构主义理论全面描写汉语〕

赵元任,1948,《北京口语语法》,李荣编译,开明书店 1952。载李荣《语文论衡》。〔最早用结构主义理论全面描写汉语〕

赵元任,1959,《语言问题》,商务印书馆 1980。

赵元任,1968,《汉语口语语法》(A Grammar of Spoken Chinese. Berkeley),商务印书馆 1979。[用结构主义方法全面描写汉语]

赵元任,1992,《中国现代语言学的开拓和发展》,清华大学出版社。[赵元任论文集]

周国光,1995,《确定配价的原则与方法》,沈阳、郑定欧《现代汉语配价语法研究》。[配价是语义层面的问题]

朱德熙,1978. 1—2,《的字结构和判断句》,《中国语文》。[提出了根据歧义指数确定向的原则]

朱德熙,1978. 3,《在黑板上写字及相关句式》,《语言教学与研究》,修订稿在《语法丛稿》。[变换分析]

朱德熙,1979,《汉语句法中的歧义现象》,《中国语文》1980. 2。[区分了隐性关系和显性关系。涉及了语义指向]

朱德熙,1979. 2,《与动词给相关的句法问题》,《方言》。[变换分析,语义特征分析]

朱德熙,1980,《现代汉语语法研究》,商务印书馆。

朱德熙,1982,《语法讲义》,商务印书馆。

朱德熙,1982. 1,《语法分析和语法体系》,《中国语文》。

朱德熙,1983. 1,《自指和转指:汉语名词化标记的、者、所、之的语法功能和语义功能》,《方言》。

朱德熙,1984. 6,《关于向心结构的定义》,《中国语文》。[引入语义组合条件限制向心结构的定义]

朱德熙,1985,《语法答问》,商务印书馆。[全面展开了对分布理论的辩护,系统展开了词组本位的理论]

朱德熙,1985. 5,《现代书面汉语里的虚化动词和名动词》,《北京大学学报》。

朱德熙,1986. 2,《变换分析中的平行性原则》,《中国语文》。

朱德熙,1990,《语法丛稿》,上海教育出版社。

后 记

从1996年开始,我在北京大学中文系开始给研究生讲授当代语言学,每年一次,每周3学时。除了在中国香港、日本工作和讲学中中断过两次,共讲授了12次。截至1996年,国内当时已经有不少学者介绍过国外语言学流派的工作,并做过很多有益的分析,学生都有机会接触到,我也受益很多。叶蜚声先生生前也多次讲过国外语言学课,也有很多当代语言学流派的介绍材料。在这样一个背景基础上,可以进一步展开对当代语言学主要理论模型的分析 and 研究。当代语言学主要理论模型并不是在讨论纯理论问题,其实这些理论模型很重视对具体语言现象的观察和分析,这需要阅读当代语言学的原典,并结合材料展开理论模型的分析,对问题的认识才会深入。这就形成了当代语言学的开课目的,所以我开课的中心不再围绕流派介绍,而主要抓住一些重要的理论模型,让学生阅读原著,然后我结合汉语和汉藏语的材料对各种理论模型进行分析、验证,对存在的问题提出自己的修正意见和解决方案。课程的重心围绕两条线展开,一条线是从结构语言学到转换语言学的转型,一条线是从同质研究到异质研究的转型。前一条线索内容多,问题复杂,主要研究共时系统问题,占据了主要的课题时间,是本书的核心内容。后一条线索也很重要,我将在思考比较成熟的时候对相关内容加以整理出版。好在北京大学中文系研究生还有一门必修课历史语言学,可以对思考同质和异质问题有帮助。

这门课需要阅读大量当代语言学经典文献,国内翻译、影印的材料已经有一部分。北京大学教务部给这门课的建设有

过几次资助,有利于我从国内外购买和复印一些文献。1998年到1999年我到香港城市大学电子工程学系王士元先生处做数理语言学访问研究,王士元先生早年也介绍研究转换生成语法,并出版过Chomsky(1957)的Syntactic Structure的最早汉译本。我在香港期间王先生给我提供了很多这方面的资料信息,并常常和我讨论转换生成语法的利弊,让我受益颇多。我的好友周铎先生多年在美国政府部门做图书信息管理工作,我能查阅到的文献他都能很快给我寄过来。另一位好友美国休斯敦医学研究中心的萧建国博士对语言学很有兴趣,非常乐意给我寄送最新成果。我要特别感谢他们在收集资料方面对我的帮助。他们置身语言学的山外,总能在黑夜中看到山上的亮光。

听课的学生从不同的角度不断提出各种问题,促使我不断展开分析和研究,争取回答学生的问題。学生的问題是推动老师研究的重要动力。不过总觉得问題没完没了,不足以成书出版。高等教育出版社袁晓波先生在打算出这本书之前和我并不相识,袁先生从李宇明教授那里了解到北京大学开课的一些情况,并认为这门课程的讲稿有出版价值。我要特别感谢两位先生对这门课程的关心。如果没有袁先生的敦促,这本教材还在继续改动中。袁先生说得有道理,改动是没有止境的,更重要的是学生上课急切需要参考书。事实上不少听课的学生也确实等了多年还没有见书出版,私下说我说话不算数,都怪我几次庄严宣布马上要出书了。

北京大学中文系邵永海老师和我的学生宋作艳、曹晋、邱立坤、杜兆金、李子鹤、许帆婷阅读了我的书稿,纠正了很多错漏。我要特别感谢他们。这本书的编辑工作难度很大,最后我要衷心感谢高等教育出版社的袁晓波先生、于晓宁先生和吴军女士为编辑本书所付出的辛勤劳作。

由结构语言学向转换生成语法转型的过程中,有很多重要

理论问题需要严格论证,本书对主要理论模型的初步分析和提出的一些观点、方法只能代表我个人不成熟的思考,很多方面都是高度探索性的,我恳切希望读者能够多提意见。

陈保亚

2009年2月20日

于北京大学畅春园

[G e n e r a l I n f o r m a t i o n]

书名 = 当代语言学

作者 = 陈保亚著

页数 = 6 2 6

SS号 = 1 2 3 0 7 8 3 8

出版日期 = 2 0 0 9 . 0 7

出版社 = 北京市：高等教育出版社

SSLIB - JPG = http : / / image5 . 5read . com / image / ss2
jpg . dll ? did = n21 & pid = CB2C0B9037A92DF65876115
F2FF54047A810D37D825D22701C916769BAEF1C7E24
C75B81061C079A3A597F38EC8E23D37A4FC28A3802F
2E6F818FFDED326595CDBC01D39C315C41EE22DE670
399DCE950E1F4C84AC213B8F7972C9E9290A2E2AC06
944B0EB71D59BCB8EF5332FE5A40C366A & jid = /

封面
版权
目录
引言

- 1 . 言语链中语素的切分
 - 1 . 1 S a u s s u r e 线条性的实质
 - 1 . 2 双项对比
 - 1 . 3 对双项对比原则的进一步限制
 - 1 . 4 对比程序和可对比程度

思考问题

- 2 . 线性序列的生成程序
 - 2 . 1 替换与连接
 - 2 . 2 扩充与扩展
 - 2 . 3 叠加

思考问题

- 3 . 单位和规则
 - 3 . 1 词作为句法单位遇到的问题
 - 3 . 2 单位及组合的可推导性
 - 3 . 3 平行周遍原则与生成规则
 - 3 . 4 平行的性质
 - 3 . 5 周遍的性质
 - 3 . 6 多个语素序列的判定
 - 3 . 7 语素库、语符库和词库
 - 3 . 8 派生、转换和语符库的关系
 - 3 . 9 短语的语类推导

思考问题

- 4 . 替换、叠加和成分线性特征
 - 4 . 1 自动机理论与自然语言的成分线性特征
 - 4 . 2 替换的有向性与有限状态语法的左线性特征
 - 4 . 3 有限状态语法的根本属性
 - 4 . 4 替换、叠加与短语结构语法
 - 4 . 4 . 1 短语结构语法 ([Σ , F]) 的基本结构
 - 4 . 4 . 2 短语结构语法的本质
 - 4 . 5 有限状态语法和语类规约模型

思考问题

- 5 . 线性原则的不充分性和转换的必要性
 - 5 . 1 H a r r i s 的线性化工作
 - 5 . 2 转换问题的提出
 - 5 . 3 C h o m s k y 关于转换存在的理由
 - 5 . 4 转换规则的必要性
 - 5 . 4 . 1 英语助动词和转换
 - 5 . 4 . 2 主动和被动的关系
 - 5 . 4 . 3 线性组合程序和不连续直接成分
 - 5 . 4 . 4 并合与转换
 - 5 . 4 . 5 非线性组合

思考问题

- 6 . 音系分析程序的非线性化
 - 6 . 1 自主音系学的不充分性
 - 6 . 2 语素基本音形与变音
 - 6 . 3 音系分析条件：语素基本音形
 - 6 . 4 从语素音位到语音规则和语素基本音形
 - 6 . 5 核心语音单位
 - 6 . 6 提取核心语音单位以语素基本音形的长度为必要条件
 - 6 . 7 提取核心语音单位的语境不能短于语素音形
 - 6 . 8 语素音形的一致性和独立语素音形的确定
 - 6 . 9 平行周遍条件与语音规则
 - 6 . 1 0 可延展性与音节的判定
 - 6 . 1 1 两类语音组合规则

- 6 . 1 1 . 1 音段组合成音节
- 6 . 1 1 . 2 音节和音节的组合
- 6 . 1 2 对立矩阵与音位的提取程序
- 6 . 1 2 . 1 对立研究的必要性
- 6 . 1 2 . 2 元音音位的提取
- 6 . 1 2 . 3 辅音音位的提取
- 6 . 1 2 . 4 音子组合数与对立环境

思考问题

- 7 . 语义组合关系的线性分析
- 7 . 1 组合规则的语义解释
- 7 . 2 次范畴化、词库和平行周遍条件
- 7 . 2 . 1 名词的次范畴化和动词的选择规则
- 7 . 2 . 2 动词次范畴化
- 7 . 2 . 3 次范畴化和平行周遍条件
- 7 . 3 深层结构及其线性特征
- 7 . 3 . 1 深层结构的存在依据
- 7 . 3 . 2 深层结构的确定
- 7 . 3 . 3 深层结构的线性顺序

思考问题

- 8 . 语义组合关系的非线性化及其功能分析
- 8 . 1 格语法的生成机制
- 8 . 2 语义格和线性深层结构的关系
- 8 . 3 语义格的初始性
- 8 . 4 词库中动词语义格的确定和直接关联原则
- 8 . 5 动词的格框架和标注
- 8 . 6 语义结构变体标注
- 8 . 7 论元结构和次范畴结构的相互独立性
- 8 . 8 文本中语义格的识别
- 8 . 9 多项式的判定问题

思考问题

- 9 . 语义结构的再次非线性化：语义表达式
- 9 . 1 深层结构在语义解释上遇到的问题
- 9 . 2 语义表达式及其深度
- 9 . 2 . 1 表达深度 1：原子命题的组合差别
- 9 . 2 . 2 表达深度 2：词汇分解
- 9 . 3 深层结构中的词汇插入问题
- 9 . 4 元语言问题
- 9 . 4 . 1 自然语言是最初始的元语言
- 9 . 4 . 2 自然语言自身解释的内部循环
- 9 . 4 . 3 获得元语言的可能性：核心语符和核心规则
- 9 . 4 . 4 核心规则的建构分析
- 9 . 5 语义表达式的初始性

思考问题

- 1 0 . 语义解释层面
- 1 0 . 1 语义格框架和投射原则
- 1 0 . 2 语义结构的语义和转换的语义
- 1 0 . 3 平行转换中的语义差别

思考问题

- 1 1 . 结构的阶与递归性
- 1 1 . 1 X 杠理论
- 1 1 . 2 限定语和补足语层次的经验依据
- 1 1 . 3 附加成分的再分类
- 1 1 . 4 两种递归方式
- 1 1 . 5 X 杠的递归性
- 1 1 . 6 短语结构的阶和语杠理论
- 1 1 . 7 各种语类的阶和杠
- 1 1 . 8 X 杠理论的核心问题

思考问题

1 2 .	句法结构关系的初始性
1 2 . 1	支配和约束
1 2 . 2	支配概念的扩展和长距离约束的传递条件
1 2 . 3	支配与结构关系
1 2 . 3 . 1	c - c o m m a n d
1 2 . 3 . 2	m - c o m m a n d
1 2 . 3 . 3	支配的本质和作用
1 2 . 4	句法结构关系的初始性

思考问题

结语：从程序到条件

符号解释

人名译名对照外文排序

人名译名对照汉语排序

主要概念译名对照外文排序

主要概念译名对照汉文排序

索引

参考书目

后记